

Agricultural Field Experiments

Analysis of Variance with GenStat®



VISIT...

LANZAROTE
Caliente.COM

Curtis Lee graduated from North Dakota State University with a degree in Agricultural Education. He completed his MS in Crop and Weed Science as a graduate research assistant in the barley breeding program.

In 1996 he started Agro-Tech, a company offering research and development services for the crop protection and seed industries. Curt operates a farm with his father Larry where they grow wheat, canola, barley, oat, pea, lentil, corn, soybeans, flax, and sunflowers.

Curt has worked as a crop consultant, vocational agriculture teacher, and research agronomist and has over 20 years of experience in the field of agricultural research and development.

Mick O'Neill graduated from the University of Sydney, specializing in Mathematical Statistics. He completed his PhD as an Assistant Lecturer, and in 1976 accepted a lectureship in Biometry in the Faculty of Agriculture, retiring in 2006 as Associate Professor.

Since 1970 he has mixed his teaching and research activities with statistical consulting to the agriculture profession, to industry and the wider community.

Mick has a keen interest in Asia, having trained several young Indonesian academics to PhD. He has conducted many short courses in Indonesia and Laos for the Australian Centre for International Agricultural Research and other organizations. Mick is a Certified GenStat Trainer.

Data sets used for this training manual are from **Statistical Procedures for Agricultural Research, Second Edition** (9780471870920 / 0471870927) by Kwanchai A. Gomez and Arturo A. Gomez. Copyright © 1984 by John Wiley & Sons, Inc. Reproduced with permission from John Wiley & Sons, Inc.

GenStat® is a registered trade mark of VSN International Ltd, Hemel Hempstead, UK

SAS® is a registered trade mark of SAS Institute Inc., SAS Campus Drive, Cary NC 27513

Contents

Purpose	2
Forward	3
History of GenStat	5
Multi-Stratum ANOVA	7
Setting up the Model	11
Experiments with no blocking	
Completely Randomized Design, single factor treatment design, equal replication	13
Completely Randomized Design, single factor treatment design, unequal replication	30
Experiments with blocking in one direction	
Randomized Complete Block Design (RCBD), single factor treatment design	31
Randomized Complete Block Design, factorial treatment design with 2 factors	41
Split-plot Randomized Complete Block Design with 2 factors	54
Strip-plot (or split-block) Design with 2 factors	61
Split-split-plot Design in an RCB layout with 3 factors	69
Strip-split-plot Design in an RCB layout with 3 factors	81
Experiments with blocking in two directions	
Latin Square	94
Further Reading	100
References	100
Appendixes	
Data Sets from Statistical Procedure for Agricultural Research	102

Purpose

Agro-Tech utilizes GenStat® software for the analysis of data from agricultural field experiments. This document is a training manual used to guide staff through a design approach to the analysis of the data using GenStat. Data sets from Gomez and Gomez (1994) are used to illustrate the analysis.

The goal of this manual is the following:

- Introduce ANOVA and limitations of this technique.
- Introduce the history and philosophy behind GenStat software.
- Understand the concept of partitioning sources of variation into stratum.
- Understand block and treatment structure and its relationship to the GenStat ANOVA procedure.
- Construct models (via the treatment and block and structure) to create ANOVA tables.
- Train staff with analysis of common treatment and experimental designs used in field plot research and ANOVA techniques via GenStat software.
- The following are covered in this manual.

Treatment designs: Single factor, two or more factors in factorial combinations

Experimental designs: CRD, RCB, Latin Square, split-plot, strip-plot, split-split plot & strip-split plot

GenStat is widely used for the analysis of data for agricultural field experiments. GenStat has a strong agricultural following in Europe, Australia and New Zealand. It is a very powerful and dependable product for designed experiments and agricultural research.

Forward

ANOVA

Modern research is concerned with the detection of small differences. This requires the use of efficient designs and methods that will most effectively reduce experimental error to obtain reliable results. In 1923, through a manure trial on potatoes, Fisher introduced the analysis of variance (ANOVA). Since then ANOVA became institutionalized as the standard of what is commonly accepted as standard statistical analysis for experimental research data. This idea remains firmly in place today (Stroup, 2013).

The analysis of variance does the following in a systematic way (Peterson, 1994).

- Partitions the total variation in the data into components associated with such sources of variation as error, treatments, and grouping categories (blocks).
- Provides an estimate of experimental error that can be used to construct interval estimates and significance tests.
- Provide a format to test the significance of several variation sources in the partition.

Sir Ronald Fisher said the analysis of variance (ANOVA) is not a mathematical theorem but rather a convenient method of arranging the arithmetic. However, it may be more insightful to understand ANOVA as a thought process of how data has arisen rather than just a process of arithmetic. Historically these computations were completed by hand calculations and the construction of the ANOVA table. Today ANOVA analysis is completed by algorithms built into the statistical software.

This manual's focus is analysis of variance using GenStat® software which will prove to be a valuable tool for the field agronomist. Knowing where the numbers come from can be

invaluable for better understanding the ANOVA table and analysis. We refer you to Gomez (1994) for a presentation of the mathematical calculations involved in the analysis of the data presented.

Limitations of ANOVA

This manual focuses on traditional ANOVA methods which have limitations. It assumes that data have approximate normality, independent observations and common variance. When it does not, a variance-stabilizing transformation has traditionally been used

In practice, common variance is often not the case and data and non-normal distributions are common. This includes spatial data and temporal data, data involving discrete, categorical or courteous response variables, multi-location and multi-year data, and repeated measures (Guber, et al, 2012, Lee, et al 2007, Stroup, 2012).

Today, more sound approaches are available for the analysis of such data. In such cases, a generalized linear mixed models approach is preferred and is available in GenStat's regression and Mixed Models menu.

Although Linear Mixed Models (LMM) are very flexible, it should be noted that it is easy to fit a wrong model and obtain misleading results with LMM. As the models becomes even more complex (Generalized Linear Mixed Models), so does the danger of misspecification of the model.

Mixed model analysis is not the objective of this manual. Nevertheless, an understanding of ANOVA and the underlying blocking and treatment structures is essential to the mixed modelling process and the transition to this type of analysis.

History of GenStat

In 1843, John Bennet Lawes founded the Rothamsted experiment station to investigate the impact of inorganic and organic fertilizers on crop yield. Lawes was an entrepreneur and scientist who founded one of the first artificial fertilizer manufacturing facilities in 1842. Lawes appointed a young chemist named Joseph Henry Gilbert and launched the first of a series of long-term field experiments. Over the next 57 years, Lawes and Gilbert established the foundations of modern scientific agriculture and the principles of crop nutrition.

Rothamsted pioneered the application of statistics in biological research when Sir Ronald Fisher was appointed in 1919 to study the accumulated results of Broadbalk, the oldest continuous agronomic experiment in the world. Fisher realized the need for improved statistical techniques over the whole range of agricultural and biological research, and the groundwork for modern applied statistics was laid by him and his colleagues during the 1920s and 1930s.

Statistical computing began at Rothamsted when Fisher's successor Frank Yates obtained an Elliot 401 computer – one of the first computers to be used away from its manufacturing base, and one of the first to be used for statistical work. This extended the tradition, started by Fisher, of conducting statistical research to solve real problems arising from biological research. The resulting new methods could now be implemented in the Rothamsted statistical programs to enable them to be used more effectively in practice. The development of GenStat at Rothamsted began in 1968, when John Nelder took over from Yates as Head of Statistics. Roger Payne took over leadership of the GenStat in 1985. Currently GenStat is being developed and marketed by VSN International (VSNi) as a spin-off company from Rothamsted and NAG

(Numerical Algorithms Group). The development group has retained its close links with the research community and Rothamsted (VSNi, 2014).

Why GenStat

The following attributes make GenStat a very useful, productive, and economical tool for the analysis of agricultural experiments.

- Provides a flexible system for analysing experimental data through its ANOVA, regression and REML facilities.
- The ANOVA is very powerful and can analyse nearly all balanced standard designs.
- Uses the concept of strata and uses the correct mean square for computing F-tests.
- It prints ANOVA tables in the conventional form you find in statistical text books.
- Contains cutting edge statistical methods that can be accessed through a menu system or a programming language.
- Has a powerful spreadsheet which allows for easy data entry, import, export and manipulation of data.
- Has effect algorithms that allow for quick creation of different experimental designs, exploration of data, analysis and printout.
- Implements good statistical practices through an intelligent menu system and comprehensive suite of diagnostic messages.
- Allows quick formation of summary tables and graphics for data visualization.
- It excels in its ability to handle multiple sources of variation which is the key to agricultural experimentation.

The development of GenStat can be traced back to the Rothamsted Experiment Station in England, one of the oldest agricultural research institutions in the world.

Multi-Stratum ANOVA

The **multi-stratum analysis of variance** is a leading principle behind the analysis agricultural data and is fundamental to understanding design itself. This tradition in design and analysis is taught at Rothamsted Research, implemented in GenStat, and utilized at Agro-Tech. A recent book, “Statistical methods in biology”, gives a detailed explanation of this and other approaches (Welham et al, 2015). This reference was used in the construction of the following training document.

In a statistical way of speaking we structure our trials into **strata** to minimize the heterogeneity of error (maximize soil uniformity) within blocks. We may further structure our trials to accommodate equipment used to apply treatments. Consequently, restrictions are imposed on layout of an experiment every time we design and conduct an experiment. These restrictions create *different structural sources of variability among the experimental units called **strata***. Each restriction in the structure of an experiment is called a **stratum**.

The **multi-stratum ANOVA** accounts for the physical structure of the experimental material or blocking imposed by the experimenter. It is an analysis approach that creates an ANOVA table with separates components for each strata defined by the structural component (block model or block structure). The variation within each stratum is partitioned into the sums of squares associated with the treatments that vary between the units at that level of the design and a residual term.

Strata are the different structural sources of variability among the experimental units.

The great advantage of the multi-stratum ANOVA is the recognition of the interplay between blocking and treatment structure so that treatment effects are always allocated to the correct strata so appropriate variance are calculated

There is an old adage in statistics, **“as the randomization is, so should the analysis be”** (Pearce, 1988). The GenStat ANOVA algorithm is true to this idea as it is a design based analysis. This is a natural approach to the analysis of data from agricultural field experiments. Very few software packages are available that create multi-stratum ANOVA tables.

A multi-stratum ANOVA table is not without limitations. It can only be formed when the explanatory and structural component obey certain conditions of balance. The simplest case of this occurs with block and treatment factors are orthogonal as in a RCB design. Although GenStat implicitly identifies terms in the structural component of the model as random, they are calculated by least square estimates as if they were fixed terms.

“The long and short of the multi-stratum ANOVA is that if you’ve specified you structure correctly then treatment terms get tested at the correct level of structure. If you are not using a multi-stratum ANOVA table or do not know how your software computes the F test’s, working out estimated means squares is essential” (S.J. Welham, personal, communication, 2015).

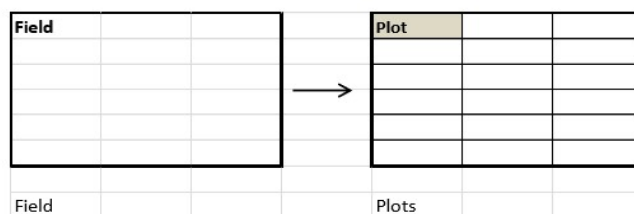
For mixed models (those with both random and fixed effects) the REML facilities should be used.

Visualizing STRATA in a Field Experiment

The reason of concerning ourselves with the concept of strata in experimental design is to accommodate and adjust for field variation and treatment application through blocking. Think of strata in terms of structural restrictions imposed on the experimental units in a field.

CRD

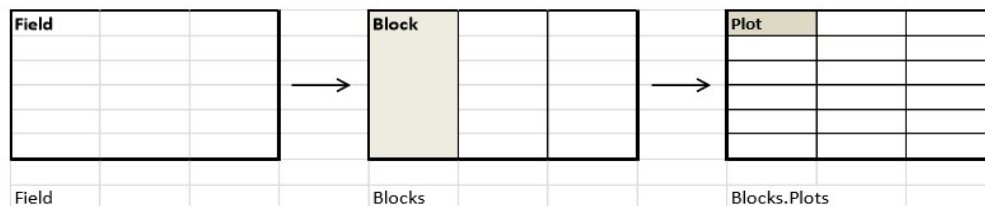
Field → Plots (this has no structural restrictions imposed)



There are **no stratum**.

RCBD and RCBD Factorial

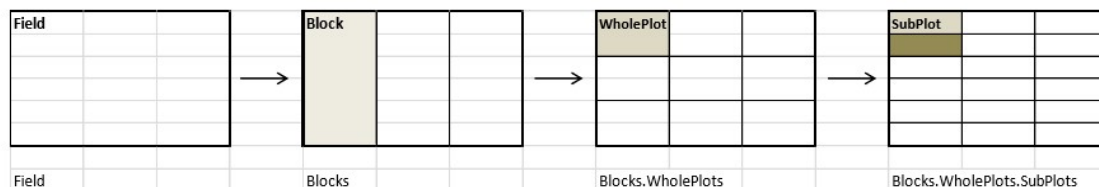
Field → Blocks → Plots



There are **two stratum**:
 1. Blocks (Variation between blocks)
 2. Blocks.Plots (Variation between plots within block)

Split Plot

Field → Blocks → WholePlots → SubPlots



There are **three stratum**:
 1. Blocks
 2. Blocks.WholePlot
 3. Blocks.WholePlots.SubPlots

Comparison of a simple ANOVA table and GenStat's Multi-Stratum ANOVA table.

Analysis of variance

Simple ANOVA Table (RCBD).

Variate: Yield

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Block	3	1944361.	648120.	5.86	
Treatment	5	1198331.	239666.	2.17	0.113
Residual	15	1658376.	110558.		
Total	23	4801068.			

A simple ANOVA does not make any distinction between describing the underlying structure of the data and those indicating the treatments applied.

Multi-Stratum ANOVA (RCBD).

Variate: Yield

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Block stratum	3	1944361.	648120.	5.86	
Block.Plot stratum					
Treatment	5	1198331.	239666.	2.17	0.113
Residual	15	1658376.	110558.		
Total	23	4801068.			

*The multi-stratum ANOVA table for the RCBD rearranges the simple ANOVA table to reflect the structure of the experiment. The RCBD has two distinct stratum, a **Block stratum** and a **Block.Plot stratum**.*

The multi-stratum ANOVA table is a general ANOVA table that preserves the distinction between the terms describing the underlying variability structure of the data (block structure) and those indicating the treatments applied (treatment structure).

Setting up the model

The GenStat model consist of two formulas. The first formula is the structural component **(block structure)** which describes the way the experimental units are established in the field before you apply the treatments.

The second formula is the explanatory component **(treatment structure)** which describes exactly what treatments are applied to the experimental units, and (possibly) sets up specific questions to be answered in the analysis, for example whether there is a linear trend in yield with increasing amounts of a fertilizer.

These structures are derived through a combination of identifiers (terms) and operators called the model formula. A model formula is a list of identifiers (name given the data structure within program) and operators defining the model terms to be analysed. The operators proved a convenient way of stating a model in a compact form. The two most common relationships between terms (factors) are **nested** and **crossed structures**.

The / **(forward slash)** operator indicates a **nested relationship**. This is a hierarchical relationship where multiple units of one structural level are entirely contain within unit at a higher level.

Block/plot = Block + Block.Plot (Blocks and plots within blocks)

The * (star) operator indicates a **crossed relationship**.

Variety * Fertilizer = Variety + Nitrogen + Variety.Fertilizer

Commonly Used Operators

Operators

Addition operator (+)	$A+B+C$ main effects of A , B , and C
Interaction operator (.)	$A.B$ interaction of A and B
Crossing operator (*)	$A*B$ is equivalent to $A+B+A.B$
Nesting operator (/)	A/B is equivalent to $A+A.B$

STRUCTURAL AND EXPLANATORY COMPONENT EXAMPLES

Example:	<u>Structural Components</u>	<u>Explanatory Component</u>
CRD:	None used	Treatment
RCBD:	Block/Plot	Treatment
Latin Square:	Row*Column	Variety
Split Plot:	Block/W_Plot/S_Plot	Variety*Nitrogen
Strip Plot:	Block/(W_Plot1*W_Plot2)	Nitrogen*Variety
Split Split Plot:	Block/W_Plot/S_Plot/SS_Plot	Nitrogen*Management*Variety
Strip-Split Plot:	Block/(Row*Column)/PlantingMethod	Variety*Nitrogen*PlantingMethod

One of Genstat's noted achievements is that it incorporated John Nelder's theory of balance into Graham Wilkinson algorithm, and pushed this concept to the limit.

In summary, it puts all the work of Fisher, Yates and Finney into a single framework so that any design can be described in terms of two formulas (Senn, 2003).

This made it possible to retain the conceptual simplicity of ANOVA type strata in the analysis, which is very intuitive for those analyzing designed experiments.

Completely randomized design

These concepts, and the way that the block and treatment structures are completed, are best described through simple examples.

Table 1.1 Data from Page 14 of Gomez and Gomez: grain yields (kg ha⁻¹) of rice resulting from different foliar and granular insecticides for the control of brown plant hoppers and stem borers, from a CRD experiment with 4 (*r*) replications and 7 (*t*) treatments

	Replicate			
Treatment	1	2	3	4
Azodrin	2387	2453	1556	2116
Control	1401	1516	1270	1077
DDT + γ -BHC	2536	2459	2827	2385
Dimecron-Boom	1997	1679	1649	1859
Dimecron-Knap	1796	1704	1904	1320
Dol-Mix (1 kg)	2537	2069	2104	1797
Dol-Mix (2 kg)	3366	2591	2211	2544

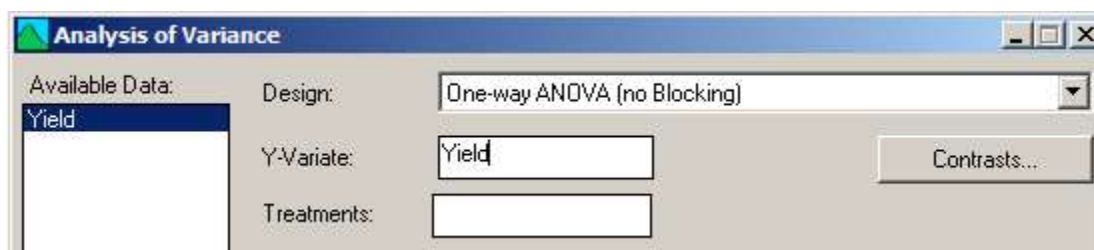
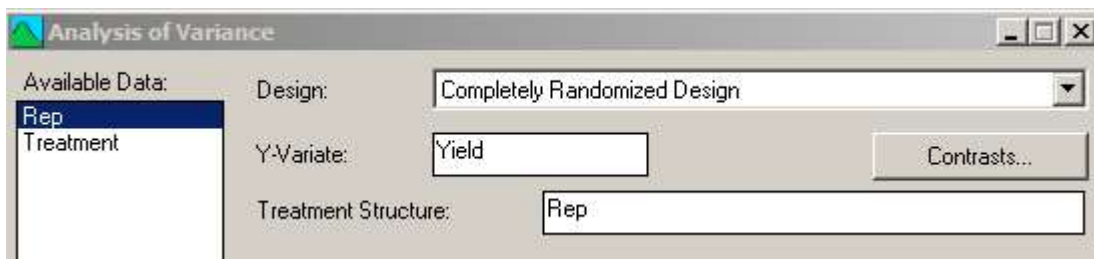
How did this data set come about? One scenario is this. The experimenters had 28 plots on which the same variety of rice was grown. Every plot had roughly the same growing conditions. Four of the 28 plots were randomly selected to have Azodrin applied, another four plots were randomly selected to be untreated (the Control treatment), and so on.

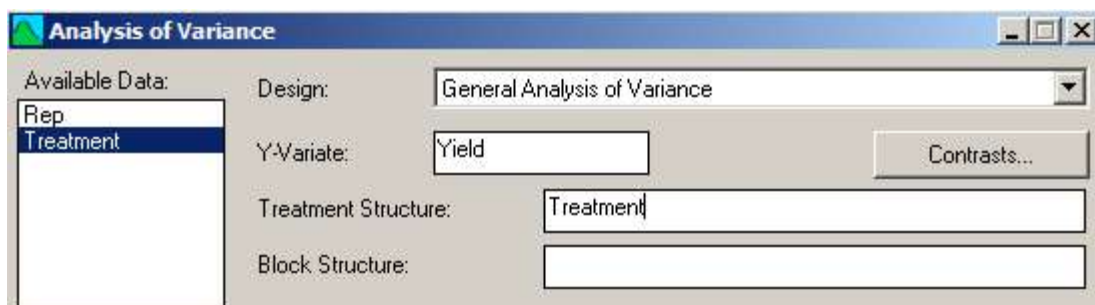
So for this simple design there is just one simple treatment structure consisting of six different insecticides plus a control. The grain yields would have to be stacked in one column (which is a *variate*, that is a column that just contains numeric data) in a GenStat spreadsheet, alongside of which would be a second column identifying which of the 7 “treatments” had been applied to the plot that produced that grain yield. This column could be named simply *Treatment* or *Insecticide*, and would have to be declared a *factor* in GenStat: this is a column that simply

contains different *levels* – in this case the 7 different insecticides (or control) applied to the rice plots. GenStat indicates a factor by placing a red exclamation mark to the left of the column name.

Row	Treatment	Yield
1	Dol-Mix (1 kg)	2537
2	Dol-Mix (1 kg)	2069
3	Dol-Mix (1 kg)	2104

GenStat has several ways of completing the analysis of this data set. In Stats > Analysis of Variance you can select Completely Randomized Design, One-way (no blocking) or General Analysis of Variance from the drop down menu alongside Design: They all produce the same analysis. All you need to enter is the response variate to be analyzed (we've named it *Yield*) and the treatment structure which is simple in this case.





The term One-way in the middle menu simply refers to the number of treatment factors in the experiment, in this case just one. There are **no blocks** in this example: the 7 treatments were allocated randomly to the 28 plots in the field, four replicates per treatment. Note that when we click inside the Y-Variate box, only potential *variates* are listed, not factors. In the Treatment Structure box, however, both *factors* and *variates* are listed because both treatments (which are factors) and covariates can be selected for analysis.

The General Analysis of Variance menu can be used for any design no matter how complex, so it is worth spending a little time to understand the various strata in an analysis.

The Block Structure in this menu was left blank simply because there are no blocks in this example. However we could have filled this in, but we need to think through the model before we discuss how.

Rice grown on the plots used in this experiment will produce some mean yield. We *expect* a treatment effect (or else why run the experiment?). That is, plots that are treated are likely to produce on average more than plots that are untreated, and whether the use of foliar or granular insecticides is better can be determined by the experimental outcome.

In summary, for this simple experiment, every yield can be expressed in a simple model:

$$\text{Yield} = \text{overall mean} + \text{individual treatment effect} + \text{error}$$

The *error* for each plot simply allows for individual yield variation. There are 28 plots, so underlying the data vector in GenStat is an error vector with 28 random error terms that allow individual plots to be slightly different from what would theoretically be expected. These errors are assumed to have the same distribution: they are all normally distributed, independent (in the sense that one plot does not affect nor is affected by another plot – can this assumption be reasonable in a field trial? – more of this later), and theoretically they all have 0 mean and a common variance which is often labelled σ^2 .

Moreover, the 28 plots are the replicates for the treatments, and, for a completely randomized design, what GenStat requires on the Block Structure is a **factor** (or **combination of factors**) that indexes through all 28 plots. This could be done in one of two ways:

1. A factor named (say) *Reps* could be inserted into the spreadsheet consisting of 28 levels and whose rows take values 1, 2, ..., 28. *Reps* would be entered as the Block Structure.
2. A factor named (say) *Rep* could be inserted into the spreadsheet consisting of 4 levels and whose rows take values {1, 2, 3, 4} repeated 7 times. Then you would combine the two factors *Rep* and *Treatment* as the Block Structure using GenStat's convention *Rep.Treatment* (1 to 4 times 1 to 7 producing the required 28 plot identifiers)

So the Block Structure in the General Analysis of Variance menu can either be left blank for a completely randomized design, or be filled with either of these two structures:

The top screenshot shows the 'Analysis of Variance' dialog box with the 'Design' field set to 'you want to capture.'. The 'Y-Variate' is 'Yield', 'Treatment Structure' is 'Treatment', and 'Block Structure' is 'Reps'.

The bottom screenshot shows the 'Analysis of Variance' dialog box with the 'Design' field set to 'General Analysis of Variance'. The 'Y-Variate' is 'Yield', 'Treatment Structure' is 'Treatment', and 'Block Structure' is 'Rep.Treatment'. The 'Available Data' list includes 'Rep', 'Reps', and 'Treatment'.

Row	Treatment	Yield	Reps	Rep
1	Dol-Mix (1 kg)	2537	1	1
2	Dol-Mix (1 kg)	2069	2	2
3	Dol-Mix (1 kg)	2104	3	3
4	Dol-Mix (1 kg)	1797	4	4
5	Dol-Mix (2 kg)	3366	5	1
6	Dol-Mix (2 kg)	2591	6	2

This demonstrates a general GenStat convention. There is only one plot shape in the field for this experiment, and hence only one stratum in the analysis of variance. We don't actually need to enter the plot structure in this example. The convention for more complex designs is this:

The final stratum can always be omitted in a GenStat Block Structure simply because GenStat will always add one in if we do not. So now we turn to the output from the Analysis of Variance of the rice yields.

Why Analysis of variance?

First, go back to the data and simply calculate the sample variance of the 28 rice grain yields.

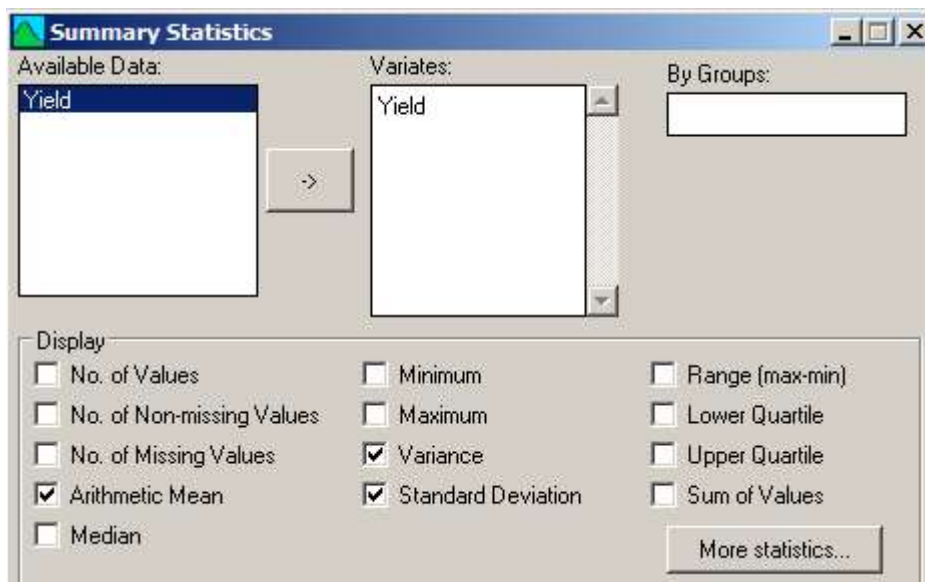
You can do this in Excel or in GenStat. In Excel, you would use the following formulae (and suppose the array of rice yields is named Yield):

=AVERAGE(Yield) for the mean

=VAR(Yield) for the usual sample variance

=STDEV(Yield) for the usual standard deviation

In GenStat, the simplest menu is found in Stats > Summary Statistics > Summary Statistics...



You would obtain a sample variance of 280,645. A sample variance of n data values has $(n-1)$ df (degrees of freedom) - because once you subtract the mean from each data value, the new values add to 0 so there is one restriction governing them. So for our 28 grain yields the sample variance has 27 df.

Now look at GenStat's default output, the first part of which is the traditional ANOVA table found on Page 16 of G&G:

Analysis of variance					
Variate: Yield					
Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Treatment	6	5587175.	931196.	9.83	<.001
Residual	21	1990238.	94773.		
Total	27	7577412.			

GenStat uses **s.s.** for Sum of Squares, **m.s.** for Mean Square and **v.r.** (variance ratio) where G&G uses Computed *F*. GenStat also provides the P value (labelled **F pr.**) used for assessing the hypothesis that all treatment *means* are equal (or, in terms of our model formulation, that all treatment *effects* are zero).

To explain further, a mean square is simply the sum of squares divided by its degrees of freedom

$$m.s. = \frac{s.s.}{df}$$

and the computed *F*, or v.r., is the ratio of the Treatment Mean Square to the Residual Mean Square

$$v.r. = \frac{\text{Treatment } m.s.}{\text{Residual } m.s.}$$

If there are no treatment effects the numerator and denominator should be roughly the same, and therefore you would not expect the variance ratio to be much larger than 1. The **P value** (F pr.) is the probability of obtaining your computed *F*, or a larger one, assuming that the treatment means are all equal. Conventionally, a 0.05 value is set for deciding when your

treatment means are significantly different: we say we are using a 5% level of significance and reject the assumption that the treatment means are all equal whenever the P value is less than 0.05. But be aware you are playing a numbers game: there are “unlucky” samples when the treatment means all equal that produce a 5% significant variance ratio – in fact 1 in every 20 times that will happen. But at the end of the day you have just the one experiment from which a decision is to be made. You are more “confident” in rejecting that the means are equal than you are in saying you’re just unlucky.

However for this experiment, there is **very strong statistical evidence** ($P < 0.001$) that the treatment means are not all equal. Note that that doesn’t imply they are all different; some are, all might be.

The ANOVA table is a conventional layout, designed in pre-computer days when hand calculation was the way the components in the ANOVA table were calculated. You’ll notice that the Treatment and Residual sums of squares add to the Total sums of squares, which means that, once the latter is calculated, only one of the former two components is required; the second can be obtained as a difference.

G&G show how to hand calculate the sums of squares; it was a valuable monograph for agricultural scientists before the advent of modern computers. However, the monograph does not mention variances, and that’s what we want to emphasize in this manual.

We saw that the sample variance was 280,645 with 27 df. The *Total* in the ANOVA table has 27 df, but its mean square is left blank (in this menu; it *is* calculated in the regression menu). Were

we to calculate this value it would (by definition) be $7,577,412/27 = 280,645$ – *which we have previously calculated*:

The Total m.s. in an ANOVA table is simply the sample variance of all the data being analyzed.

Hence the name Analysis of Variance: it is a process of looking at the variance of the data (ignoring any structure like treatments and blocks), then breaking it down into components which are also variances that can be explained, and using those components to form a decision in terms of the aims of the experiment.

So how are the other two component mean squares related to variances?

Look firstly at the treatment means. GenStat's Summary Tables menu is like Excel's Pivot Table procedure. In GenStat, simply select the Variate to be reported on and the factor (*Group:*) for which individual statistics are required – in this case treatments:

Treatment	Mean	Variance
Azodrin	2128	166678
Control	1316	35487
DDT + γ -BHC	2552	37473
Dicecron-Boom	1796	26556
Dicecron-Knap	1681	64601
Dol-Mix (1 kg)	2127	93631
Dol-Mix (2 kg)	2678	238986

There is quite a bit of variation among the 7 treatment means. The control mean is the lowest, presumably because without any treatment the rice has suffered from the plant hoppers and stem borers. The means then range to a maximum of 2678 kg ha^{-1} , with Dol-Mix (2 kg) apparently the best of the 6 insecticides.

Each of the treatment means is based on 4 plot yields. We could calculate the *variance* of the 7 treatment means to measure their variability – its value is 232,799. The variance of all the data was based on 28 individual plots; the variance in treatment means is based on seven means each of 4 plots. So to put the latter on the same unit plot basis, we multiply the sample variance of the treatment means by 4, the number of plots that each is based on, and the answer is $4 \times 232,799 = 931,196$ – which is the *Treatment m.s.* Hence:

The Treatment m.s. in an ANOVA table is simply the sample variance of all the treatment means being analyzed, scaled up by the number of plots each mean is based on.

Keeping on this tack, the sample variance of the 4 plots that received Azodrin is 166,678 and this is based on $4 - 1 = 3$ df. This would be a possible estimate of experimental error (the parameter for which we defined as σ^2). But the same argument would hold for every treatment, assuming a common variance: each is a potential estimate of σ^2 , and each has 3 df. So a better estimate is to *average* these 7 variances: the average of 232,799 to 238,986 is 94,773 and the df for this estimate is the combined df of the individual estimates, in this case $7 \times (4 - 1) = 21$. This is the *Residual m.s.* in the ANOVA table. Hence:

The Residual m.s. in an ANOVA table for an experiment with no blocks is simply the average of the sample variances of the individual treatments.

It can be shown that the *Residual m.s.* is also the sample variance of the **residuals** from the model, provided the *Residual df* are used in the calculation of the variance formula, and that this statement is a general result for any design.

What are the fitted values?

Each model leads to a set of fitted values, which are simply the sample estimates of the parameters in the model. For our CRD model

$$\text{Yield} = \text{overall mean} + \text{individual treatment effect} + \text{error}$$

the *overall mean* is estimated by the mean of all the data (2040 kg ha⁻¹) and the treatment effects are estimated as the differences between the individual treatment means and the overall mean. So for this design, the fitted value for a given yield is the estimate of (*overall mean + individual treatment effect*), which simplifies to *the individual treatment mean*.

For example, the fitted values for each of the 4 plots on which Azodrin was applied is 2128 kg ha⁻¹.

What are the residuals?

The residuals are the differences between the observed yields and the fitted yields. For example, the four individual plot yields on which Azodrin was applied are 2387, 2453, 1556, 2116, and hence the four residuals are 2387-2128 = 259, 2453-2128 = 325, 1556-2128 = -572, 2116-2128 = -12. Notice that these add to 0, as will be the case for the residuals from the other treatments.

Remember that for normal data 95% of all random samples will fall within ± 2 (roughly) standard deviations of the mean. The theoretical errors are assumed to be normal with 0 means, and consequently GenStat flags any residual outside of this bound; it prints the value of the residual and its estimated standard error:

Message: the following units have large residuals.

units 5	688.	s.e. 267.
units 15	-572.	s.e. 267.

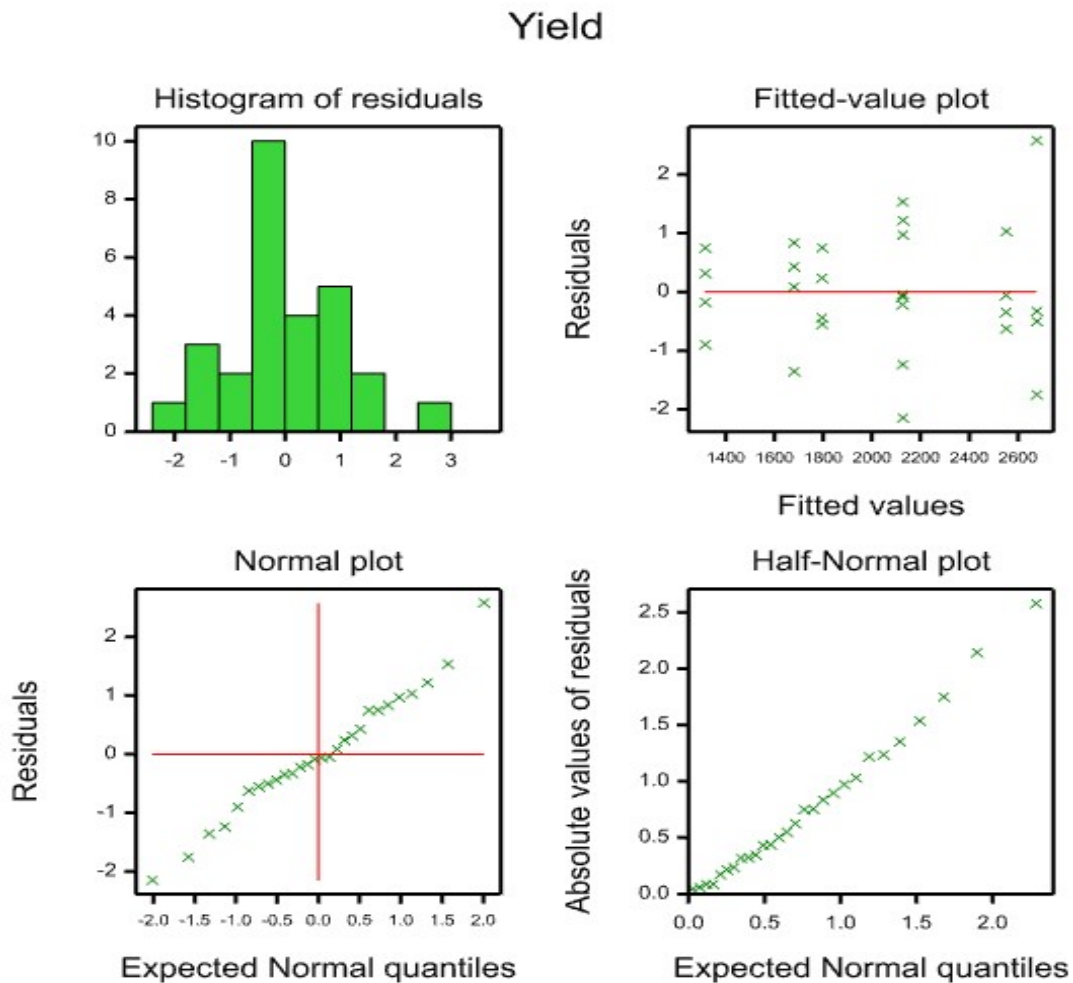
You should not be too disturbed at these messages, unless the ratio of the residual to its standard error is large, say more than 3, which is unlikely (it would happen for roughly only 3 of every 1,000 data points by chance). With 28 yields in this experiment you would expect 1 or two residuals to be flagged.

What is more important is to examine residuals in field order when that is applicable, because there should be no pattern in the residuals in sign or size. Clusters of large positive residuals might indicate a high fertility area that was not accounted for in setting up the experiment (and hence in the analysis). Nor should there be trends when residuals are plotted against fitted values: if larger fitted values have larger residuals (a plot of residuals might appear to fan out), that would indicate a problem with your assumption about equal variance. It may be that a $\log(\text{yield})$ variate should be analyzed instead, or a newer analysis used that allows for changing variance.

The kind of experiment that we are analyzing here is typical of a possible change in variance: control plots (which here are untreated) all show low yields due to the damage by pests, so you might expect the control variance to be different to the variance for treated plots. Indeed, the latter may also vary with the success of the insecticide. This problem, however, is not for this manual.

For the record, the residual plot is available once the analysis is run by clicking Further Output > Residual Plots. Our preference is to also select Standardized so it becomes easy to identify how

many and which points are outside the $(-2, +2)$ bounds. This plot does suggest that treatments with smaller fitted values (so the Control and the two treatments with low levels of the insecticide) may be behaving differently with respect to variance, but with only 4 replicates each the evidence for this is not strong.



GenStat's default output includes means and standard errors of means. Optionally you can request *least significant differences (l.s.d. values)*:

Tables of means

Variate: Yield

Grand mean 2040.

Treatment	Azodrin 2128.	Control 1316.	DDT + γ -BHC 2552.	Dimecron-Boom 1796.
Treatment	Dimecron-Knap 1681.	Dol-Mix (1 kg) 2127.	Dol-Mix (2 kg) 2678.	

Standard errors of differences of means

Table	Treatment
rep.	4
d.f.	21
s.e.d.	217.7

Least significant differences of means (5% level)

Table	Treatment
rep.	4
d.f.	21
l.s.d.	452.7

The l.s.d. values are used in two complementary ways.

Firstly, any two means that differ by at least the l.s.d. value are significant at 5% at least. So the Control is significantly different to the Dimecron-Boom insecticide treatment since $1316 - 1796 = 480 > 452.7$; and, since the other means are larger than the latter, the control is significantly different (again, at 5% at least) to *every* insecticide treatment.

Secondly, the l.s.d. value is what you add and subtract to the difference of two means to obtain a (95%) confidence interval (CI) for that “true” difference. So a 95% CI for the difference in mean yield for say Dimecron-Boom and Dimecron-Knap is $1796 - 1681 \pm 452.7$ which leads to an interval $(-338, 568) \text{ kg ha}^{-1}$. This interval includes 0 and says that the two treatments are not statistically significant (at 5%).

Contrasts

While discussing differences in treatment means we can make use of GenStat's powerful contrasts facility in the Analysis of Variance menu. We won't look at what are called orthogonal contrasts (in GenStat known as Regression type) but at simple contrasts which simply obey the rule that the coefficients of a contrast add to 0.

A brief background.

When we compared the Control (call its mean say $\bar{Y}_1$) and Dimecron-Boom (say $\bar{Y}_2$) we essentially took the mean difference ($\bar{Y}_1 - \bar{Y}_2$), which can be thought of as $(+1)\bar{Y}_1 + (-1)\bar{Y}_2$. The coefficients of $\bar{Y}_1$ and $\bar{Y}_2$ in this linear function (+1 and -1) add to 0.

Suppose next you believed that treatment 2 and treatment 3 were very similar, and wanted to see if treatment 1 differed from them (on average). We base our decision on the mean difference

$$\left(\bar{Y}_1 - \frac{\bar{Y}_2 + \bar{Y}_3}{2} \right)$$

which can be re-expressed as $\left(+1 \times \bar{Y}_1 - \frac{1}{2} \times \bar{Y}_2 - \frac{1}{2} \times \bar{Y}_3 \right)$ and the set of coefficients

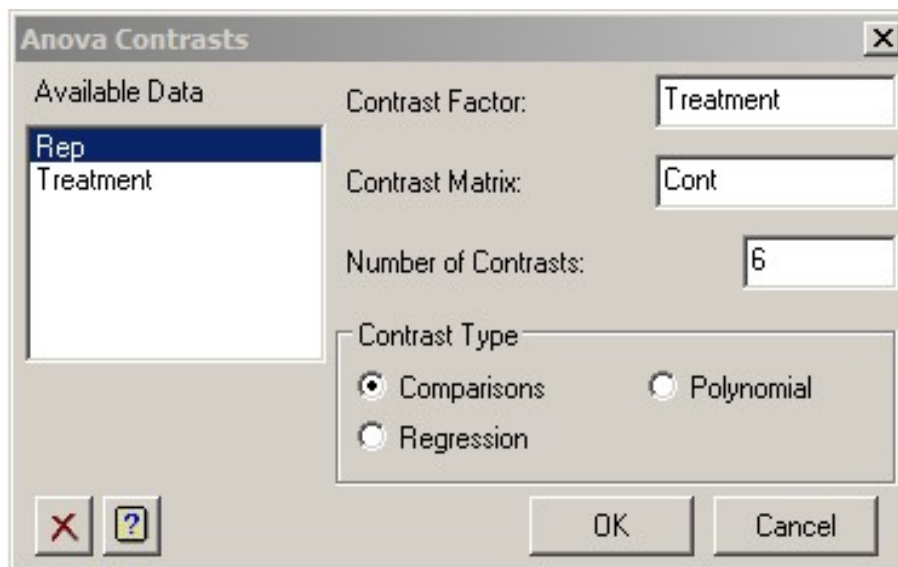
$\left(+1, -\frac{1}{2}, -\frac{1}{2} \right)$ still add to 1. Fractions are sometimes awkward to enter (think of 1/3); in this example we can multiply by 2 and use (2, -1, -1) instead: it makes no difference to the F tests.

If you think of the 7 means in our example being arranged in a list $\{\bar{Y}_1, \bar{Y}_2, \dots, \bar{Y}_7\}$ there are many such comparisons you might make.

One natural question (contrast) to consider might be whether applying an insecticide *or not* is beneficial to rice production. This would be a contrast of the Control with all 6 insecticides, and, generalizing the last contrast, would be (6, -1, -1, -1, -1, -1, -1). The signs don't matter: $\bar{Y}_1 - \bar{Y}_2$ answers the same questions as does $\bar{Y}_2 - \bar{Y}_1$. They just need to be consistent so we could use (-6, 1, 1, 1, 1, 1, 1) just as well.

Another contrast might be whether Dol-Mix (2 kg) produces better yields than Dol-Mix (1 kg). If these are the final two treatments in the list then the contrast would be (0, 0, 0, 0, 0, 1, -1), the 0s indicating that those treatments are not considered in this particular question.

To do this in GenStat, either highlight the *Treatment* factor and click the Contrasts button, or just click the Contrasts button and select the *Treatment* factor. Then indicate the number of contrasts you wish to make – let's ask 6 different questions:



A table pops up waiting for you to indicate the coefficients for the 6 questions you're interested in, for example:

1. Control vs Treated (-6, 1, 1, 1, 1, 1, 1)
2. Dimecron-Boom vs Dimecron-Knap
3. Azodrin vs the two Dimecron treatments
4. Dol-Mix (1 kg) vs Dol-Mix (2 kg)
5. The two Dimecron treatments vs the two Dol-Mix treatments
6. DDT + γ -BHC vs Dol-Mix (2 kg)

Row	Rows	Azodrin	Control	DDT + γ -BHC	Dimecron-Boom	Dimecron-Knap	Dol-Mix (1 kg)	Dol-Mix (2 kg)
1	Cont vs Treat	1	-6	1	1	1	1	1
2	Dim-M vs Dim-K	0	0	0	1	-1	0	0
3	Az vs Dim	-2	0	0	1	1	0	0
4	Dol 1 vs Dol 2	0	0	0	0	0	-1	1
5	Dim vs Dol	0	0	0	1	1	-1	-1
6	DDT vs Dol 2	0	0	-1	0	0	0	1

Complete the table of contrast coefficients and click back into the Analysis of Variance menu.

The ANOVA table has the overall test of all seven treatment means, followed by an individual F test and P value for each of the comparisons among the seven treatments you decided to make:

Analysis of variance

Variate: Yield

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Treatment	6	5587175.	931196.	9.83	<.001
Cont vs Treat	1	2443742.	2443742.	25.79	<.001
Dim-M vs Dim-K	1	26450.	26450.	0.28	0.603
Az vs Dim	1	404561.	404561.	4.27	0.051
Dol 1 vs Dol 2	1	607753.	607753.	6.41	0.019
Dim vs Dol	1	1762920.	1762920.	18.60	<.001
DDT vs Dol 2	1	31878.	31878.	0.34	0.568
Residual	21	1990238.	94773.		
Total	27	7577412.			

Completely randomized design with unequal replication

Unequal replication for a simple treatment structure causes no real difficulty. The formulae discussed in the previous section are slightly modified (so replicates are used as weights inside each formula). Otherwise the concepts are the same.

Where the treatment structure is more complex and is unequal replication, the analysis is more complex and will be left for a later discussion.

An example of a more complex treatment design might be where we have one factor, *herbicide* (with different herbicides forming the levels), in combination with a second factor, *rate of application* (with different rates for each herbicide). This is known as a *factorial treatment structure*. *Herbicide × Rate of Application* is called a two-way treatment structure, *Herbicide × Rate of Application × Time of Application* is a three-way treatment structure, and so on.

In GenStat a shortcut way for defining a factorial treatment structure is simply to list the factors with an * between them, so *Herbicide*Rate* and *Herbicide*Rate*Time* and so on.

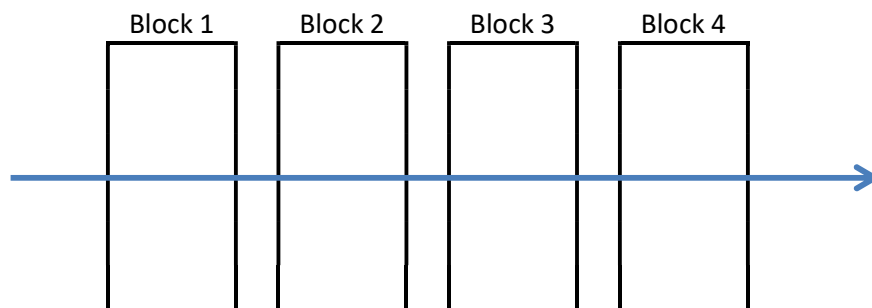
The example in G&G on Page 19 is almost an example of a three-way treatment design, but not all combinations are present and there are other treatments (such as the Control) that are not in combination with anything. In G&G the experiment is used with all the treatment combinations set out in a single treatment simply to illustrate the unequal replication feature of the analysis.

Randomized Complete Block Design (RCBD)+simple treatment structure

As G&G explain, the RCBD has been the most common design used in agricultural research. The reason it is selected is usually because it is often not possible to find sufficiently many plots that are all alike in growing conditions prior to the randomization of treatments to plots.

We won't repeat the discussion on Pages 20-22 of G&G but the discussion is excellent. What we concentrate on is the simple step from the field layout to the GenStat analysis.

Suppose you have a fertility gradient of some sort or another **left to right**. We would need, therefore, to form (say 4) blocks *perpendicular* to that gradient:

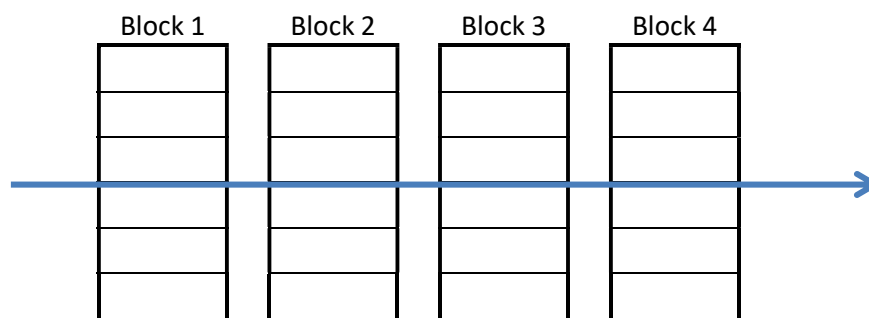


Immediately, then, we have a stratum which we can name the **block stratum**. Note, however, that the land in Block 1 differs in growing conditions to the land in every other block because of the fertility trend moving left to right. So the large area we call Block 1 is unreplicated anywhere else in the research area. Blocks are unreplicated; we need replication of a factor in order to form a test of that factor (for example, in the first example we had 4 replicates of every treatment and were able to test whether the treatment means were all equal). Hence, technically, blocks cannot be tested in an analysis of variance.

Another point to make before moving on is the fact that whatever conclusions we draw from this experiment in which we formed blocks, we would like to generalize our findings to other areas in which the growing conditions are similar to those in our experimental area. In that sense we regard the blocks we just formed as a *random example* of the types of land we want to make recommendations about. We call the *Block* factor a *random factor* so we can make generalizations, unlike a *fixed factor* such as a *Rate* of seeding which might differ from 25 kg/ha to 150 kg/ha in steps of 25 kg/ha (so 6 levels). In the latter case we confine our attention to (within) that range of seeding, and not comment on say 200 kg/ha.

The next step in the design of our experiment is to apply the treatments to the blocks. The G&G example on Page 26 has 4 blocks and the 6 seeding rates just defined. So in each block, we need to construct 6 similar plots, then randomize the 6 treatments to those 6 plots, block by block:

Forming plots in each block:



Randomizing treatments (rate of seeding) in each block:

Block 1	Block 2	Block 3	Block 4
125	75	100	75
50	125	50	25
150	100	125	100
25	25	75	150
75	150	150	50
100	50	25	125



In forming the 6 plots in each block we have constructed a *second stratum*, a plot in a block, whose shape is just $\frac{1}{6}$ th of the block shape. We can write what we did pictorially:

1. First form blocks in the field, thereby generating a factor *Block* with 4 levels.
2. Next take each block and divide it into 6 plots, thereby forming a *Plot* factor with 6 levels. Physically the process of forming plots in each block can be represented as *Block/Plot*.

This is exactly what GenStat allows you to use in the General Analysis of Variance for an RCB design. However think through the two strata we constructed in the field:

-  There are 4 blocks, so the *Block* stratum is simply the factor named *Block*.
-  There are $4 \times 6 = 24$ plots altogether, so to index every plot for the *Plot* stratum we need to use the construct discussed in the first example, *Block.Plot*.

So in GenStat, *Block/Plot* is simply a shortcut for *Block + Block.Plot*.

Since GenStat allows the smallest (last) stratum to be omitted, we need only use *Block* in the General Analysis of Variance for an RCB design. However you should keep in mind that the full

structure is *Block/Plot = Block + Block.Plot* because with the newer REML analysis you might need to use the full structure for various reasons.

It is worth thinking through what happens when the treatments are compared. The yield from Block 1 associated with 25 kg/ha seeding rate is made up of an overall mean, plus or minus a component associated with that seeding rate, plus or minus a component due to the fact that this plot is in a block on the left of the fertility spectrum. In fact any yield in the experiment can be expressed as:

$$\text{Yield} = \text{overall mean} + \text{block effect} + \text{treatment effect} + \text{error}$$

which is similar to the CRD model but contains the additional block effect.

So the mean yield for the 25 kg/ha treatment includes one yield from every block, so on average contains an average block effect. But so does every other treatment mean. When you *compare* two treatment means, the average block effect, which is present in both means, disappears in the treatment mean difference. So every treatment mean difference is an estimate of that real treatment effect, if there is one.

This will not be the case if every treatment does not occur in each block, as happens for example when one plot is ruined by say a flood in the corner of the research field. To make the point, suppose that, in the highest fertile block, the plot that had treatment 1 randomized to it was accidentally destroyed. A comparison between the mean of treatment 1 with any other mean would now be most unfair, because the other treatment mean has a yield coming from the highest fertile block and treatment 1 does not.

This might be easier to see using the model and taking expectations. Let's use τ_i for the i^{th} treatment effect and β_j for the j^{th} block effect. We want to estimate $(\tau_1 - \tau_2)$, remembering that treatment 1 is missing from the highly fertile block 1:

$$E(\bar{Y}_1 - \bar{Y}_2) = (\tau_1 - \tau_2) + \frac{\beta_2 + \beta_3 + \beta_4}{4} - \frac{\beta_1 + \beta_2 + \beta_3 + \beta_4}{5}$$

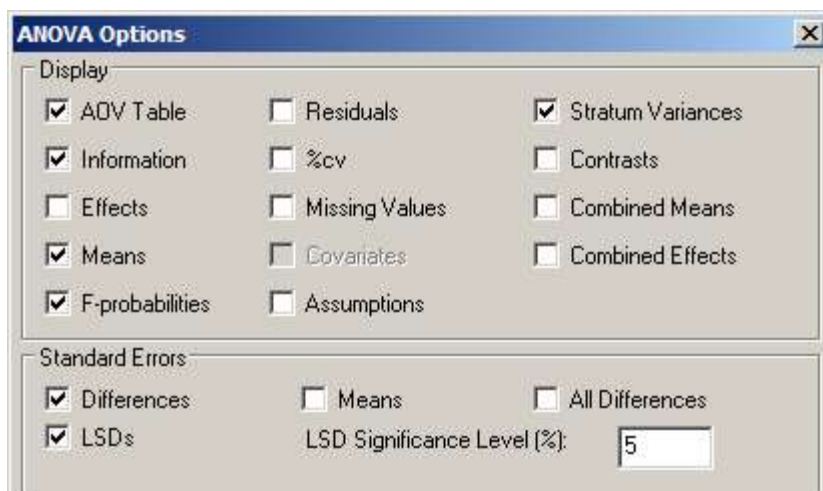
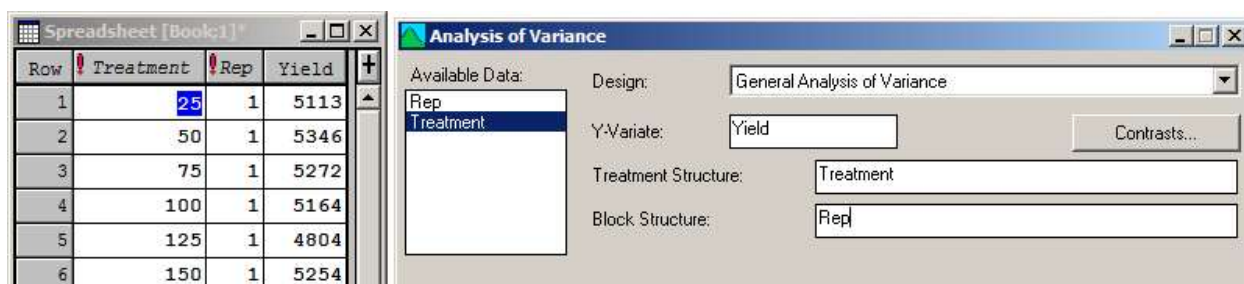
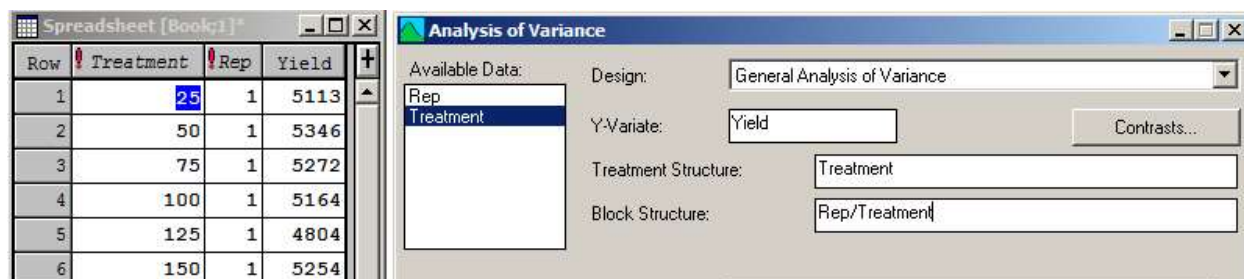
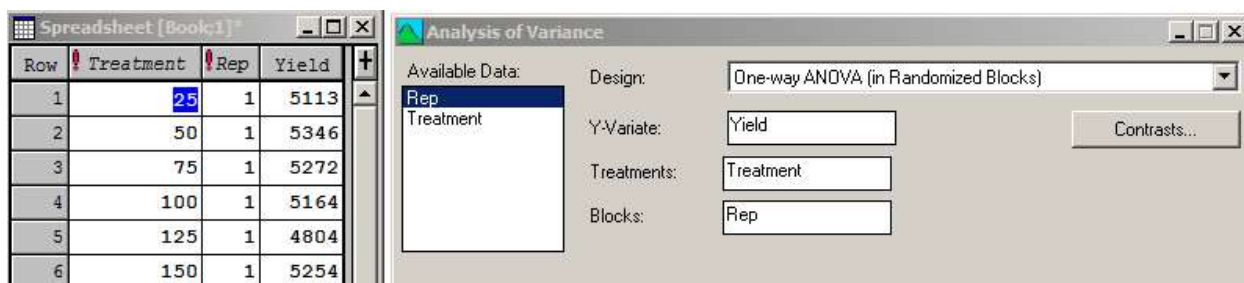
You can see that the block effects no longer cancel out in this difference. So for missing data, we need a different analysis. The **C** in RCB indicates a *complete* set of treatments in each block

RCB data from G&G Page 26. Grain yield, kg/ha:

Rate, kg seed/ha	Block				Treatment mean
	1	2	3	4	
25	5,113	5,398	5,307	4,678	5,124
50	5,346	5,952	4,719	4,264	5,070
75	5,272	5,713	5,483	4,749	5,304
100	5,164	4,831	4,986	4,410	4,848
125	4,804	4,848	4,432	4,748	4,708
150	5,254	4,542	4,919	4,098	4,703
Block mean	5,159	5,214	4,974	4,491	

Notice that the yields are low towards the right of the field, so the decision to block was possibly justified. By blocking with 4 blocks, 3 df have been lost from the CRD residual term in the ANOVA, so there is slightly less precision for comparing means across seeding rates; but a possible large gain by controlling the variation in yields.

In GenStat we can again use the specialist or general analysis of variance menu. Blocks are called Rep in the data set. The next three menus all give the same ANOVA. The only difference between the second and third menus is in the way GenStat names the plot stratum (it uses Rep.*Units* stratum when the plots within a block are dropped from the Block Structure):



In addition to the means and the standard error of a treatment mean difference (sed) that GenStat prints as defaults, in Options you can request LSD values and Stratum Variances, as well as the standard error of a mean (in case you want to add this common value as an error bar in an Excel scatter plot).

The sample variance of the 24 yields is 208,742 and this has 23 df. Once again, we should expect to see a **Total s.s.** of $23 \times 208,742 = 4,801,068$ in the ANOVA.

Analysis of variance

Variate: Yield

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Rep stratum	3	1944361.	648120.	5.86	
Rep.Treatment stratum					
Treatment	5	1198331.	239666.	2.17	0.113
Residual	15	1658376.	110558.		
Total	23	4801068.			

The sample variance of the treatment means (which are each based on 4 yields) is 59,917 and, using the same logic as with the CRD ANOVA, $4 \times 59,917 = 239,666$ will be the **Treatment m.s.**

The sample variance of the block means (which are each based on 6 yields) is 108,020 and, using the same logic as with the CRD ANOVA, $6 \times 108,020 = 648,120$ will be the **Block m.s.**

We have not yet discussed the analysis of a factorial treatment design, but when we do you will see that the *Residual* is the interaction between *treatments* and *blocks*. We don't expect that treatments will behave differently in each block, and hence we use this term in the construction of the F statistic (or variance ratio, v.r.) for comparing treatment means.

Note that GenStat, like G&G, does *not* construct an F test for blocks. Blocks are in a stratum on their own. There is no replication of blocks, so as a fixed effect blocks are untestable. Almost every other statistical package incorrectly provides a P value for blocks.

So from the ANOVA table we conclude there is insufficient statistical evidence ($P = 0.113$) that not all means are equal. That's not to say that every pairwise comparison will be not significant. There are only 4 replicates of each treatment, and even if there are differences the overall F test is not powerful enough to detect an overall difference.

The 5% least significant difference (LSD) value can be used in two ways.

Firstly, any two treatment means that differ in absolute value by more than the LSD value are in fact significantly different at a 5% level. With an LSD value of 501 it would appear that a 75kg/ha seeding rate is superior to 150 kg/ha since this difference is $5304 - 4703 = 601$ kg/ha.

Secondly, a 95% confidence interval for a mean difference is obtained by adding and subtracting the LSD value from the difference. So the 75kg/ha seeding rate is superior to 150 kg/ha by 601 kg/ha, although the true difference is likely to be within the interval 601 ± 501 , or (100 kg/ha, 1,102 kg/ha).

Tables of means

Variate: Yield

Grand mean 4960.

Treatment	25	50	75	100	125	150
	5124.	5070.	5304.	4848.	4708.	4703.

Standard errors of means

Table	Treatment
rep.	4
d.f.	15
e.s.e.	166.3

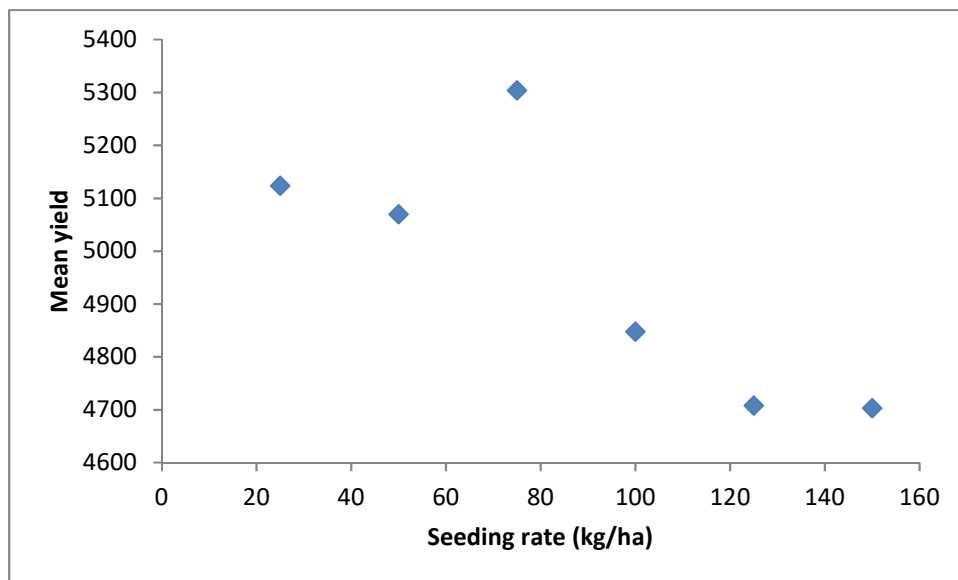
Standard errors of differences of means

Table	Treatment
rep.	4
d.f.	15
s.e.d.	235.1

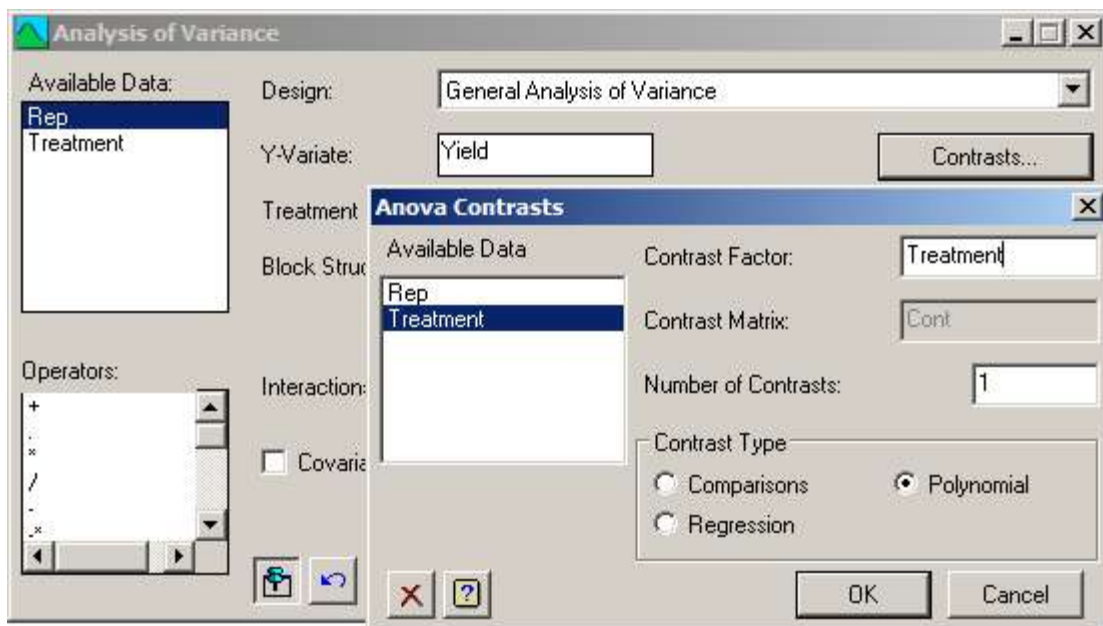
Least significant differences of means (5% level)

Table	Treatment
rep.	4
d.f.	15
l.s.d.	501.1

When the means are plotted the 75 kg/ha rate appears unusual. We do not know enough of the background to ask why this is the case. But if we take it at face value, we might wonder whether there is a linear decline in mean yield with increasing seedling rate.



Again, we can select Contrasts in the ANOVA menu, this time selecting Polynomial and setting the degree of the polynomial to 1 (linear):



Analysis of variance

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Rep stratum	3	1944361.	648120.	5.86	
Rep.Treatment stratum					
Treatment	5	1198331.	239666.	2.17	0.113
Lin	1	760035.	760035.	6.87	0.019
Deviations	4	438296.	109574.	0.99	0.442
Residual	15	1658376.	110558.		
Total	23	4801068.			

There is a significant linear decrease ($P=0.019$). Notice that of the variability among treatment means that can be explained, a simple linear model explains $760,035/1,198,331 = 63\%$.

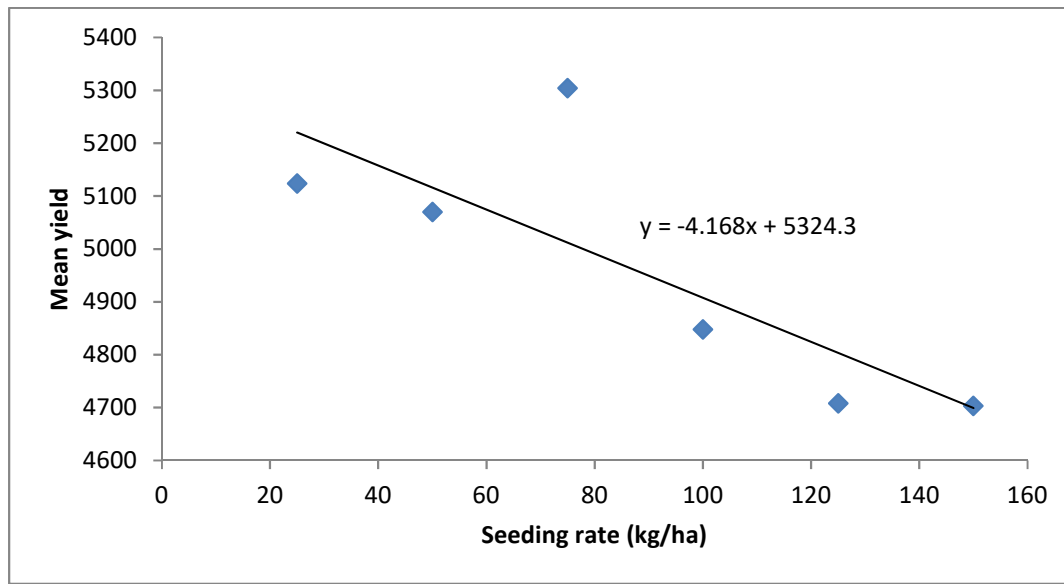
If you also request **Contrasts** in Options, the *slope* of the linear regression is printed out:

Treatment contrasts

Lin -4.2, s.e. 1.59, ss.div. 43750.

The yield declines by 4.2 kg/ha for every 1 kg/ha additional seeding rate. The mean seeding rate is 87.5, and the actual regression line is

$$\text{Yield} = \text{mean yield} - 4.2 \times (\text{Rate} - \text{mean rate}) = 4,960 - 4.2 (\text{rate} - 87.5) = 5,324 - 4.2 \times \text{Rate:}$$



Randomized Complete Block Design + factorial treatment structure

Making the treatment structure more complex does not cause any difficulty in the field provided that all treatments are completely randomized in every block. When a more complex allocation of treatments to field plots occurs, that is when a more complex analysis is required.

The following example from Page 92 of G&G illustrates a two-factor treatment design set out in 4 randomized blocks. You'll see there is no pattern in the allocation of treatments to plots in each block. There are 3 varieties ($V_1 = 6966$, $V_2 = P1215936$, $V_3 = \text{Milfor } 6(2)$), each in combination with 5 levels of a nitrogen fertilizer (N_0 to $N_4 = 0, 40, 70, 100, 130$ kg/ha).

Rep	Treatment layout					Grain yield, t/ha				
1	V ₃ N ₂	V ₂ N ₁	V ₁ N ₄	V ₁ N ₁	V ₂ N ₃	5.822	4.956	5.874	4.788	5.664
	V ₃ N ₀	V ₁ N ₃	V ₃ N ₄	V ₁ N ₂	V ₃ N ₃	4.192	6.034	5.864	4.576	5.888
	V ₂ N ₄	V ₃ N ₁	V ₂ N ₀	V ₁ N ₀	V ₂ N ₂	5.458	5.25	2.846	3.852	5.928
2	V ₂ N ₃	V ₃ N ₃	V ₁ N ₁	V ₂ N ₀	V ₂ N ₁	5.362	5.524	4.936	3.794	5.128
	V ₁ N ₃	V ₃ N ₂	V ₁ N ₂	V ₁ N ₄	V ₂ N ₄	5.276	4.848	4.454	5.916	5.546
	V ₁ N ₀	V ₃ N ₄	V ₂ N ₂	V ₃ N ₁	V ₃ N ₀	2.606	6.264	5.698	4.582	3.754
3	V ₁ N ₁	V ₃ N ₀	V ₁ N ₀	V ₃ N ₁	V ₁ N ₄	4.562	3.738	3.144	4.896	5.984
	V ₂ N ₂	V ₁ N ₂	V ₁ N ₃	V ₂ N ₄	V ₃ N ₄	5.81	4.884	5.906	5.786	6.056
	V ₂ N ₀	V ₃ N ₂	V ₂ N ₁	V ₂ N ₃	V ₃ N ₃	4.108	5.678	4.15	6.458	6.042
4	V ₁ N ₂	V ₂ N ₂	V ₂ N ₄	V ₁ N ₀	V ₂ N ₀	3.924	4.308	5.932	2.894	3.444
	V ₁ N ₃	V ₃ N ₁	V ₁ N ₄	V ₁ N ₁	V ₂ N ₃	5.652	4.286	5.518	4.608	5.474
	V ₃ N ₀	V ₂ N ₁	V ₃ N ₂	V ₃ N ₃	V ₃ N ₄	3.428	4.99	4.932	4.756	5.362

So, in the field, each block has to consist of 15 plots that are all assumed to be homogeneous.

The design shown here, however, *may* be compromised. If the blocks are (fairly) contiguous, then any fertility trends or differences down from block 1 to block 4 may well be reflected in the three rows of each block as well. Examining residuals in field order might show up this inner trend: we will do this post-analysis on the assumption that this was the field layout for this experiment (it is unclear from G&G if this was the case).

Strictly speaking, in GenStat the Block Structure would be Block/Plot as with the previous example: firstly blocks are constructed, and then 15 plots within each block were formed. The Plot factor would have 15 levels containing the numbers 1 to 15, or they could contain labels such as the combination of V and N applied to each of the plots, as shown here:

Spreadsheet [Book1]*

Row	row	col	Block	V	N	VN	Plot	Yield
1	1	0	1	3	2	13	V3N2	5.822
2	1	1	1	2	1	7	V2N1	4.956
3	1	2	1	1	4	5	V1N4	5.874
4	1	3	1	1	1	2	V1N1	
5	1	4	1	2	3	9	V2N3	
6	2	0	1	3	0	11	V3N0	
7	2	1	1	1	3	4	V1N3	
8	2	2	1	3	4	15	V3N4	
9	2	3	1	1	2	3	V1N2	
10	2	4	1	3	3	14	V3N3	
11	3	0	1	2	4	10	V2N4	
12	3	1	1	3	1	12	V3N1	
13	3	2	1	2	0	6	V2N0	
14	3	3	1	1	0	1	V1N0	
15	3	4	1	2	2	8	V2N2	
16	1	0	2	2	3	9	V2N3	
17	1	1	2	3	3	14	V3N3	

Analysis of Variance

Available Data: Block, N, N_No, Plot, V, VN, VN_code, V_no

Design: General Analysis of Variance

Y-Variate: Yield

Treatment Structure: V*N

Block Structure: Block/Plot

Operators: +, -, *, /, *

Interactions: All interactions.

☐ Covariates

Run Cancel

Alternatively, the Block Structure could simply be Block since the final stratum can always be omitted in GenStat (as explained previously, adding Plot simply allows that factor to be named in the smaller stratum of the ANOVA).

We could treat this experiment as a one-way treatment design with 15 treatments. We have a factor called VN in the spreadsheet that would allow this analysis:

Analysis of variance

Variate: Yield

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Block stratum	3	2.5998	0.8666	5.73	
Block.Plot stratum					
VN	14	44.5783	3.1842	21.05	<.001
Residual	42	6.3528	0.1513		
Total	59	53.5309			

There is a strong treatment effect, but you would have to look very carefully at the 15 means to make any sense of the differences. Instead, from this analysis we really like:

an idea of whether any change in yield with increasing N is consistent for every variety; and if it is consistent (the trends will be parallel), then:

is there a difference in mean yield among the varieties, averaged over nitrogen levels and

is there a change in mean yield with increasing N, averaged over varieties?

Notice that we look at the interaction *before* we interpret average (or main) effects. We only interpret the main effects when the interaction is not significant.

The reason for this is simple enough: if there is an interaction, then the optimal level of N for one variety is likely to be different to that for another variety. So your recommendation to a farmer will be provisional on which variety is being sown. If there is no interaction, the recommendation will be consistent for all varieties and that is the main effect for varieties.

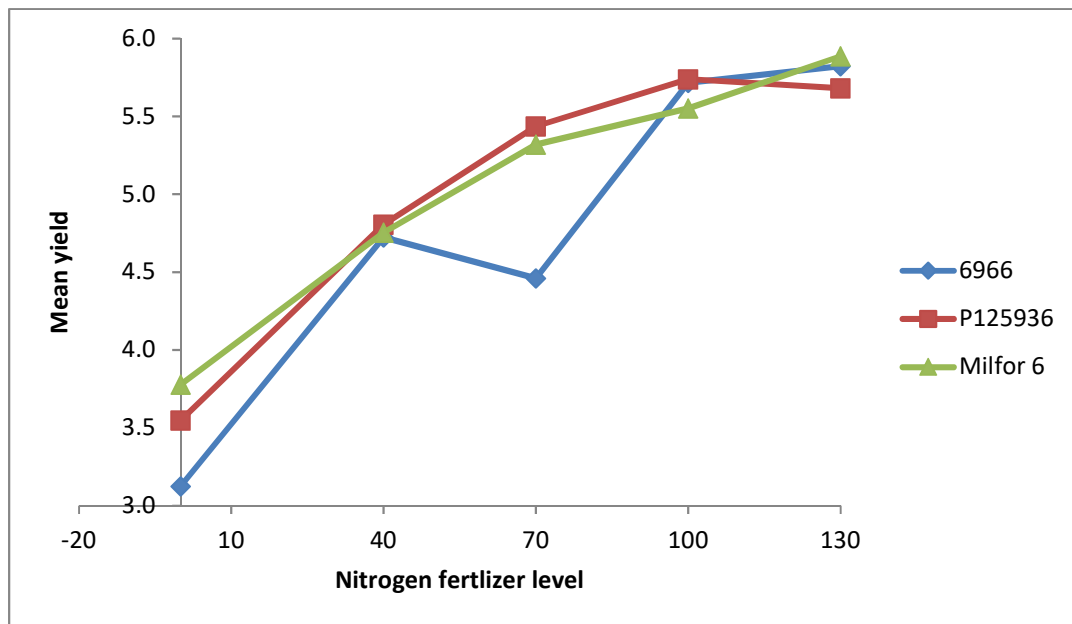
The two way means (t/ha) are as follows

Variety	Nitrogen fertilizer (kg/ha)					Variety means
	0	40	70	100	130	
6966	3.124	4.723	4.459	5.717	5.823	4.769
P125936	3.548	4.806	5.436	5.739	5.680	5.042
Milfor 6	3.778	4.753	5.320	5.553	5.886	5.058
Nitrogen means	3.483	4.761	5.072	5.670	5.797	

There appears to be a definite increase in mean yield with increasing N across the varieties.

Whether the means of the three varieties are significantly different is at this stage. And to get a

sense of whether the increase in nitrogen means is consistent across the three varieties a plot is more helpful:



There appears to be some similarity in the trends, with variety 6966 behaving oddly at 70 kg/ha N.

So, how are these questions resolved as part of the ANOVA? We use the principles already outlined.

To test whether there is a variety effect, we look at the three variety means 4.769, 5.042 and 5.058. Each mean is an average of $4 \times 5 = 20$ plots in the field: each of 4 blocks contributes 5 plots that had a particular variety (just a different level of N). So we would expect that the ANOVA would have a component **Variety m.s.** equal to:

$20 \times \text{var}(4.769, 5.042, 5.058) = 20 \times 0.026 = 0.526$ (see the ANOVA table on the next page)

To test whether there is a nitrogen effect, we look at the five nitrogen means 3.483, 4.761, 5.072, 5.670 and 5.797. Each mean is an average of $4 \times 3 = 12$ plots in the field: each of 4 blocks contributes 3 plots that had a particular nitrogen level (just a different variety). So we would expect that the ANOVA would have a component **Nitrogen m.s.** equal to:

$$12 \times \text{var}(3.483, 4.761, 5.072, 5.670, 5.797) = 12 \times 0.859 = 10.309 \text{ (see the ANOVA table)}$$

Calculation of the interaction m.s. begins to get complex and is best left to statistical software.

Basically, we mentioned that for a simple RCBD, the Residual m.s. is simply a Block $\times$ Treatment interaction. So technically, for a table of two-way means with v varieties and n nitrogen levels, we would obtain the **residuals** for each of vn means in the two-way table by subtracting the particular variety mean and the particular nitrogen mean from each two-way mean. Then we would calculate the sample variance of these values, dividing not by $(vn-1)$ but by

$$(vn-1) - (v-1) - (n-1) = (vn-v-n+1) = (v-1)(n-1)$$

These are the interaction degrees of freedom: note that the formula is a product of the degrees of freedom of the main effects making up the interaction. This simple formula allows us to know the df for any complex interaction.

The ANOVA from GenStat

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Block stratum	3	2.5998	0.8666	5.73	
Block.combo stratum					
V	2	1.0528	0.5264	3.48	0.040
N	4	41.2347	10.3087	68.15	<.001
V.N	8	2.2907	0.2863	1.89	0.087
Residual	42	6.3528	0.1513		
Total	59	53.5309			

Message: the following units have large residuals.

Block 1 combo V2N0	-0.878	s.e.	0.325
Block 3 combo V2N1	-0.846	s.e.	0.325
Block 4 combo V2N2	-0.805	s.e.	0.325

Tables of means

Variate: Yield

Grand mean 4.956

V	6966	P125936	Milfor 6			
	4.769	5.042	5.058			
N	0	40	70	100	130	
	3.483	4.761	5.072	5.670	5.797	
V	N	0	40	70	100	130
6966		3.124	4.723	4.459	5.717	5.823
P125936		3.548	4.806	5.436	5.739	5.680
Milfor 6		3.778	4.753	5.320	5.553	5.886

Standard errors of differences of means

Table	V	N	V
			N
rep.	20	12	4
d.f.	42	42	42
s.e.d.	0.1230	0.1588	0.2750

Least significant differences of means (5% level)

Table	V	N	V
			N
rep.	20	12	4
d.f.	42	42	42
l.s.d.	0.2482	0.3204	0.5550

So the ANOVA table confirms:

- there is insufficient statistical evidence to suggest that the increase in means with increasing N is not the same for all three varieties ($P=0.087$) – this is the interaction V.N;
- there is strong statistical evidence that, averaged across varieties, the means for different levels of N are statistically different ($P<0.001$) – this is main effect of N;
- there is enough statistical evidence that, averaged across nitrogen levels, the means for different varieties are statistically different ($P=0.040$) – this is main effect of V. The *lsd* values indicate that variety 6966 differs from both P125936 and Milfor 6.

Of course, had interest in the comparison of variety 6966 with each of the others been of interest prior to the analysis this could have been incorporated into the ANOVA. Similarly, the trend in mean yield across levels of N could be explored by say setting the Polynomial level to 2 (for quadratic).

So in the General Analysis of Variance menu,

- highlight N in V*N, select Contrasts, select Polynomial and change the Number of Contrasts to 2. Then
- highlight V in V*N, select Contrasts, select Comparisons and ensure the Number of Contrasts is 2. Then define and name the matrix of contrasts, as with the last example.

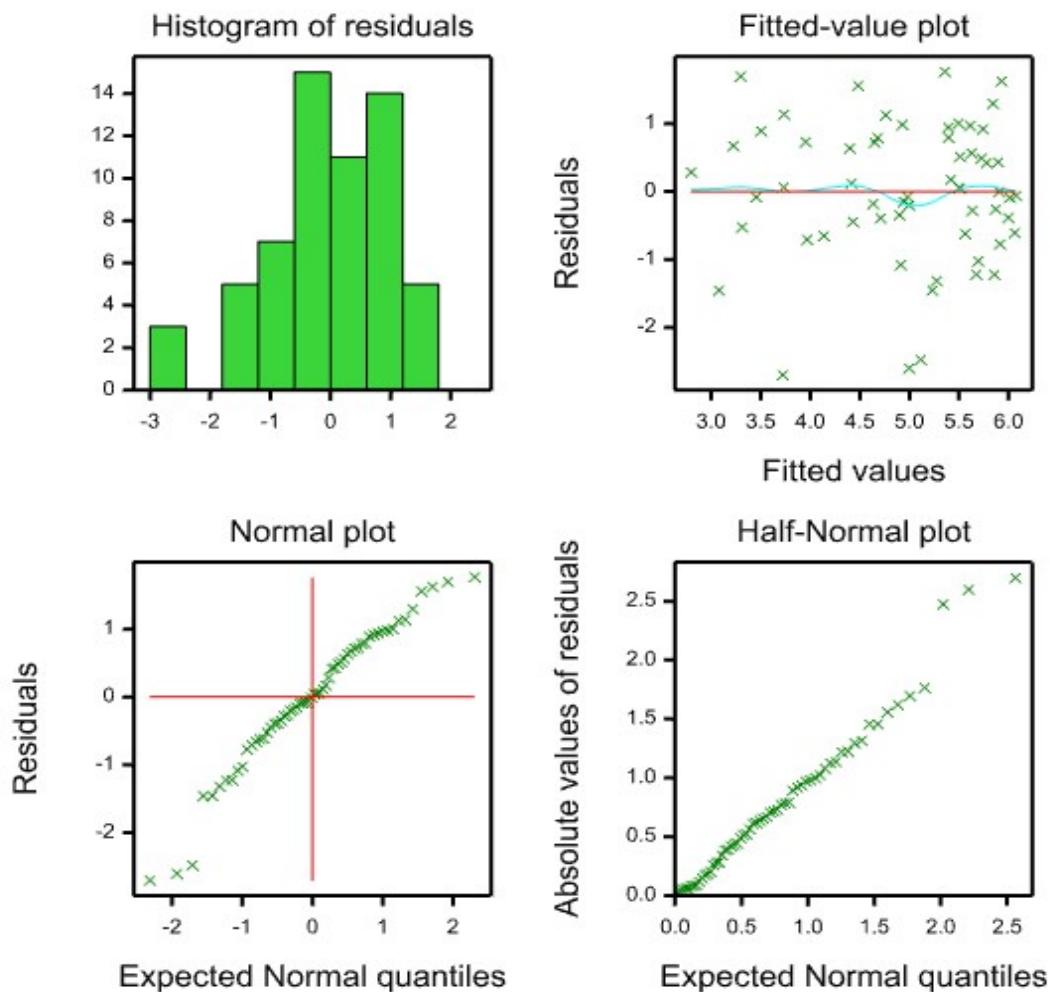
Row	 _Rows_	6966	P125936	Milfor 6
1	6966_P125	1	-1	0
2	6966_Mil	1	0	-1

GenStat automatically changes the Treatment Structure to COMP(V;2;Cont)*POL(N;2) and gives this ANOVA instead:

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Block stratum	3	2.5998	0.8666	5.73	
Block.combo stratum					
V	2	1.0528	0.5264	3.48	0.040
6966_P125	1	0.7431	0.7431	4.91	0.032
6966_Mil	1	0.8335	0.8335	5.51	0.024
N	4	41.2347	10.3087	68.15	<.001
Lin	1	36.7679	36.7679	243.08	<.001
Quad	1	3.4779	3.4779	22.99	<.001
Deviations	2	0.9890	0.4945	3.27	0.048
V.N	8	2.2907	0.2863	1.89	0.087
6966_P125.Lin	1	0.2846	0.2846	1.88	0.177
6966_Mil.Lin	1	0.3784	0.3784	2.50	0.121
6966_P125.Quad	1	0.3193	0.3193	2.11	0.154
6966_Mil.Quad	1	0.0033	0.0033	0.02	0.884
Residual	42	6.3528	0.1513		
Total	59	53.5309			

The actual P values for the variety contrasts are now given (0.032 for the first and 0.024 for the second). The trend in N is complex: not simply linear ($P < 0.001$) or quadratic ($P < 0.001$). The Deviations represents what has not been modelled: so in this case combines power 3 and power 4 components. (Through 5 points a power 4 curve fits exactly.) Biologically, a simple power model is not a good representation of the effect of N.

Yield



Finally (though this should have been done up front), we can check residuals. Click Further Output, Residual Plots and Standardized (so the bulk of residuals lie between -2 and +2). The scatter appears random, as expected (see the plot on the previous page). There are three residuals smaller than -2 (and these were flagged in to ANOVA output); these correspond to plots that had variety 2 planted in them. Three out of 60 residuals outside (-2, +2) is exactly 5%, what is expected to happen just by chance.

Notice that GenStat offered Contour Plots and a table of residuals displayed in the Output window in field layout. To do this, the data should be sorted to more easily set up an X and Y coordinate for each plot. You need to imagine the experiment layout superimposed on an X-Y axis. Then the top left hand corner plot in the field (in row 1 and column of block 1) would have a Y value of 12 (since there are $3 \times 4 = 12$ rows in the 4 blocks) and an X value of 1. The plot in Block 4, row 3 column 1 would represent coordinates Y=1, X=1. Using shading to identify the blocks, the full coordinate system of the layout in the field is as follows:

Y=12	V3N2	V2N1	V1N4	V1N1	V2N3
Y=11	V3N0	V1N3	V3N4	V1N2	V3N3
Y=10	V2N4	V3N1	V2N0	V1N0	V2N2
Y=9	V2N3	V3N3	V1N1	V2N0	V2N1
Y=8	V1N3	V3N2	V1N2	V1N4	V2N4
Y=7	V1N0	V3N4	V2N2	V3N1	V3N0
Y=6	V1N1	V3N0	V1N0	V3N1	V1N4
Y=5	V2N2	V1N2	V1N3	V2N4	V3N4
Y=4	V2N0	V3N2	V2N1	V2N3	V3N3
Y=3	V1N2	V2N2	V2N4	V1N0	V2N0
Y=2	V1N3	V3N1	V1N4	V1N1	V2N3
Y=1	V3N0	V2N1	V3N2	V3N3	V3N4
(0,0)	X=1	X=2	X=3	X=4	X=5

So in the spreadsheet, we have sorted so that the 12 data values in column 1 all appear in order. Then we pointed to the labelled row and right clicked, selected Fill, started the process at 12, ended at 1 and had to have an increment of -1 in order to descend 12, 11, ..., 1:

Row	row	col	Block	V	N	VN	combo	Yield
1	12	1	1	Milfo	2	13	V3N2	5.822
2	11	1	1	Milfo	0	11	V3N0	4.192
3	10	1	1	P1259	4	10	V2N4	5.458
4	9	1	2	P1259	3	9	V2N3	5.362
5	8	1	2	6966	3	4	V1N3	5.276
6	7	1	2	6966	0	1	V1N0	2.606
7	6	1	3	6966	1	2	V1N1	4.562
8	5	1	3	P1259	2	8	V2N2	5.81
9	4	1	3	P1259	0	6	V2N0	4.108
10	3	1	4	6966	2	3	V1N2	3.924
11	2	1	4	6966	3	4	V1N3	5.652
12	1	1	4	Milfo	0	11	V3N0	3.428
13	12	2	1	P1259	1	7	V2N1	4.956
14	11	2	1	6966	3	4	V1N3	6.034
15	10	2	1	Milfo	1	12	V3N1	5.25
16	9	2	2	Milfo	3	14	V3N3	5.524
17	8	2	2	Milfo	2	13	V3N2	4.848
18	7	2	2	Milfo	4	15	V3N4	6.264
19	6	2	3	Milfo	0	11	V3N0	3.738
20	5	2	3	6966	2	3	V1N2	4.884

Fill Column with a Numerical Sequence

row

Sequence

Starting Value:

12

Ending Value:

1

Increment:

-1

Number of Repeats:

1

☐ Copy Down existing values over missing
 ☐ Ignore restricted/filtered rows
 ☐ Fill Selected Rows only
 ☐ Fill all Columns in Selection

Fill from cell

☒ Top
 ☐ Current

Fill to

☒ Bottom
 ☐ End of List

Current cell to fill from:

4

OK

Apply

Cancel

Help

Preview

12

11

10

9

8

7

6

5

4

3

2

1

12

11

10

9

8

The two coordinate columns must be left as variates. So back in Residual Plots, identify the X and Y coordinates and select Final stratum only residuals (so that the block effect is also removed)

ANOVA Further Output

Display

☐ AOV Table
 ☐ Residuals
 ☐ Stratum Variances
 ☐ Information
 ☐ %cv
 ☐ Contrasts
 ☐ Effects
 ☐ Missing Values
 ☐ Combined Means
 ☐ Means
 ☐ Covariates
 ☐ Combined Effects
 ☒ F-probability
 ☐ Assumption tests

Standard Errors

☐ Differences
 ☐ Means
 ☐ All Differences
 ☐ LSDs

LSD Significance Level: 5

Graphics

Residual Plots...

Means Plots...

Power Calculations...

Permutation Test...

ANOVA Residual Plots

Available Data:

Y

Yield

col

row

Type of Plot

☒ Fitted Values
 ☒ Half Normal
 ☒ Normal
 ☒ Histogram
 ☐ Added Variable

Variable:

Type of residual:

☐ Simple
 ☒ Standardized

Residuals in Field Layout

☒ Contour plot
 ☐ Shade plot
 ☒ Display Table

X-Coordinates: col

Y-Coordinates: row

Method: Final stratum only

Run

Cancel

52 | Page

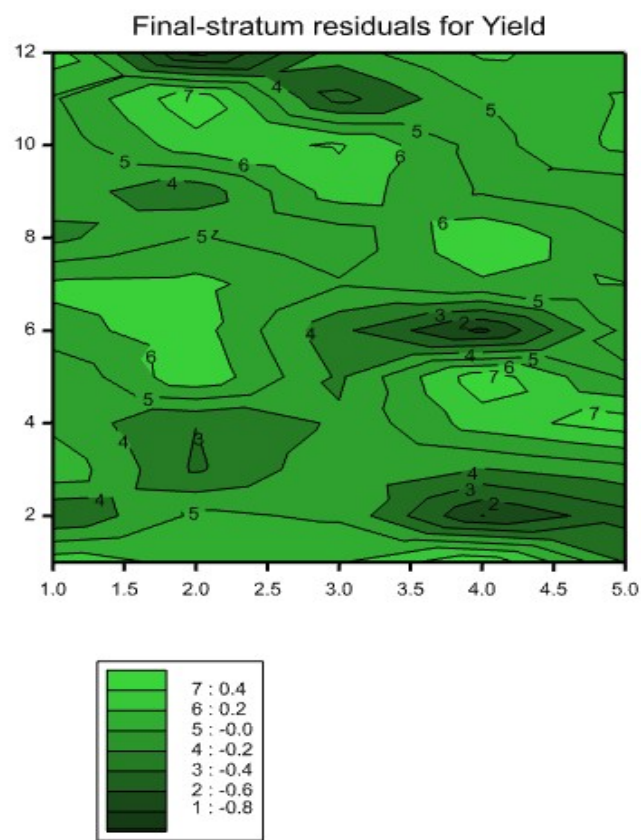
©Agro-Tech, Inc. and Statistical Advisory & Training Services Pty Ltd

The residuals in correct field order are displayed in the Output window and a contour plot is generated.

Final_stratum_residuals					
col	1	2	3	4	5
row					
12	0.3257	-0.8783	-0.0905	0.2342	0.0926
11	-0.0263	0.5517	-0.4740	-0.0013	0.2186
10	-0.1253	0.3157	0.4215	-0.0848	0.2576
9	-0.1118	-0.3335	0.3060	-0.0208	-0.1449
8	-0.2518	0.0155	-0.1275	0.3697	0.0176
7	0.2377	0.2565	0.0200	0.1677	0.2071
6	0.1407	0.2900	-0.3518	-0.8463	0.0571
5	-0.1988	0.3660	-0.2303	0.5282	-0.0274
4	-0.0598	-0.3970	-0.1703	0.2992	0.5066
3	0.1592	-0.4280	-0.0478	-0.2129	-0.0654
2	-0.3988	0.0385	-0.0293	-0.8054	-0.4739
1	0.3202	0.1370	0.1837	0.5741	-0.2019

Of course there should be no pattern in size or sign throughout the field. The contour plot is a graphic way of identifying patches of large or small residuals. There is evidence of a patchy nature in growing conditions in this field.

There are designs that allow for fertility trends in more than one direction. There are also modern methods of analysis which attempt to model more complex trends (see the REML manual for some of these).



Split-plot randomized complete block design

The next design is useful when there are two or more treatment factors and the total number of treatment combinations is too large to achieve a simple RCB layout. The design is a three-stage one.

1. Form **blocks** in the field;
2. In each block, form large areas (called **main-plots** or **whole-plots**) to which one of the treatment factors (called the whole-plot treatment) is randomized into;
3. In each main plot, form **sub-plots** or **split-plots** to which the second treatment (called the split-plot treatment) is applied to.

So there are three strata in this experiment. The Treatment Structure is an extension of that for a simple RCBD, paralleling what was done in the field: **block/whole-plot/sub-plot**. Note that either or both of the whole-plot and split-plot treatments could themselves be factorial combinations of factors. We'll illustrate though using G&G simple two-factor example from Page 102. The formation of whole-plots described in G&G could be problematic for the same reasons as in the previous example: if there really is some fertility difference among the blocks, then any difference could well be reflected within each block as well. We will, however, take the example on face value: the blocks may well not be contiguous, and the 6 main-plots may well have similar growing conditions in a given block.

The example has 6 levels of a nitrogen fertilizer (N_0 to $N_6 = 0, 60, 90, 110, 150, 180$ kg N/ha) randomized into the six whole-plots in each of three blocks. In each whole-plot, 4 varieties ($V_1 = R8, V_2 = IR5, V_3 = C4-63, V_4 = Peta$) were randomized into the 4 split-plots. The actual field layout is not provided so the data are simply presented for analysis.

whole-plot 1
whole-plot 2
whole-plot 3
whole-plot 4
whole-plot 5
whole-plot 6

Block 1

whole-plot 1
whole-plot 2
whole-plot 3
whole-plot 4
whole-plot 5
whole-plot 6

Block 2

whole-plot 1
whole-plot 2
whole-plot 3
whole-plot 4
whole-plot 5
whole-plot 6

Block 3

Grain yield (kg/ha), Page 102 of G&G

Variety	Block 1	Block 2	Block 3
0 kg N/ha			
C4-63	3464	2944	3142
IR5	3944	5314	3660
IR8	4430	4478	3850
Peta	4126	4482	4836
60 kg N/ha			
C4-63	4768	6004	5556
IR5	6502	5858	5586
IR8	5418	5166	6432
Peta	5192	4604	4652
90 kg N/ha			
C4-63	6244	5724	6014
IR5	6008	6127	6642
IR8	6076	6420	6704
Peta	4546	5744	4146
120 kg N/ha			
C4-63	5792	5880	6370
IR5	7139	6982	6564
IR8	6462	7056	6680
Peta	2774	5036	3638
150 kg N/ha			
C4-63	7080	6662	6320
IR5	7682	6594	6576
IR8	7290	7848	7552
Peta	1414	1960	2766
180 kg N/ha			
C4-63	5594	7122	5480
IR5	6228	7387	6006
IR8	8452	8832	8818
Peta	2248	1380	2014

Replicates of treatments are what are used in the construction of F tests for that treatment. We have pointed out that blocks are unreplicated, so they form a stratum on their own and we do not test for blocks.

The 6 nitrogen fertilizers were randomized into the whole-plots in each block, the application being made uniformly across each whole-plot. However, each Nitrogen mean from a single whole-plot is a mean of the 4 split-plot yields within that whole-plot. The set of means is

Nitrogen	Block 1	Block 2	Block 3	Nitrogen means
0	3991.0	4304.5	3872.0	4055.8
60	5470.0	5408.0	5556.5	5478.2
90	5718.5	6003.8	5876.5	5866.3
120	5541.8	6238.5	5813.0	5864.4
150	5866.5	5766.0	5803.5	5812.0
180	5630.5	6180.3	5579.5	5796.8
Block means	5369.7	5650.2	5416.8	

So each of the overall Nitrogen means is based on $4 \times 3 = 12$ replicates and hence the **Nitrogen m.s.** is $12 \times \text{var}(4055.8, 5478.2, 5866.3, 5864.4, 5812, 5796.8) = 12 \times 507,153 = 6,085,840$.

Similarly, each of the 3 block means is based on $4 \times 6 = 24$ plot split-yields, so

Block m.s. = $24 \times \text{var}(5369.7, 5650.2, 5416.8) = 24 \times 22,554 = 541,288$.

To test the effect of the nitrogen fertilizer we construct a **Residual m.s.** based on an RCB ANOVA of these whole-plot two-way means.

Next, varieties were randomized into each of the 6 whole-plots in each block. Hence each variety has $3 \text{ block} \times 6 \text{ split-plot} = 18$ yields from split-plots. The variety means are 5564.4,

6155.5, 6553.6, 3642.1 so

Variety m.s. = $18 \times \text{var}(5564.4, 6155.5, 6553.6, 3642.1) = 18 \times 1,664,594 = 29,962,700$.

The split-plot yields are used to obtain a **Residual m.s.** (different from that based on the larger whole-plots) used for testing for varieties.

Next, the interaction is calculated from the two-way table of Variety $\times$ Nitrogen means. Each mean is based on 3 split-plot means, one from each block. So the same **Residual m.s.** used to test for Variety is used to test the Variety.Nitrogen interaction.

So the organization of the split-plot RCB ANOVA reflects the way the Blocks, whole-plots and split-plots were formed:

The **Total m.s.** is the variance of all 72 data values, as usual, and turns out to be 2,883,773.

Component	df	m.s.
Block stratum		
Block	$(3-1) = 2$	541,288
Block.Variety stratum		
Nitrogen	$(6-1) = 3$	6,085,840
Residual (1)	$(3-1) \times (6-1) = 10$	To be calculated
Block.Variety.Nitrogen stratum		
Variety	$(4-1) = 5$	29,962,700
Nitrogen.Variety	$(6-1) \times (4-1) = 15$	To be calculated
Residual (2)	36	To be calculated
Total	$(72-1) = 71$	2,883,773

Residual (1) is the **Res m.s.** from the whole-plot replicates and Residual (2) is the **Res m.s.** from the split-plot replicates. We have entered the degrees of freedom for the latter as a *difference*

(71 minus the rest), but it is an average of six RCBD errors, one for each whole plot in an analysis of blocks and nitrogen, hence the df are $6 \times (3-1) \times (4-1) = 36$.

The full ANOVA is given over the next two pages.

Analysis of variance

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Rep stratum	2	1082577.	541288.	3.81	
Rep.Nitrogen stratum					
Nitrogen	5	30429200.	6085840.	42.87	<.001
Residual	10	1419679.	141968.	0.41	
Rep.Nitrogen.Variety stratum					
Variety	3	89888101.	29962700.	85.71	<.001
Nitrogen.Variety	15	69343487.	4622899.	13.22	<.001
Residual	36	12584873.	349580.		
Total	71	204747916.			

Tables of means

Grand mean 5479.

Nitrogen	0	60	90	120	150	180
	4056.	5478.	5866.	5864.	5812.	5797.
Variety	C4-63	IR5	IR8	Peta		
	5564.	6156.	6554.	3642.		
Nitrogen	Variety	C4-63	IR5	IR8	Peta	
0		3183.	4306.	4253.	4481.	
60		5443.	5982.	5672.	4816.	
90		5994.	6259.	6400.	4812.	
120		6014.	6895.	6733.	3816.	
150		6687.	6951.	7563.	2047.	
180		6065.	6540.	8701.	1881.	

Standard errors of differences of means

Table	Nitrogen	Variety	Nitrogen Variety
rep.	12	18	3
s.e.d.	153.8	197.1	445.5
d.f.	10	36	43.53
Except when comparing means with the same level(s) of			
Nitrogen			482.8
d.f.			36

Least significant differences of means (5% level)

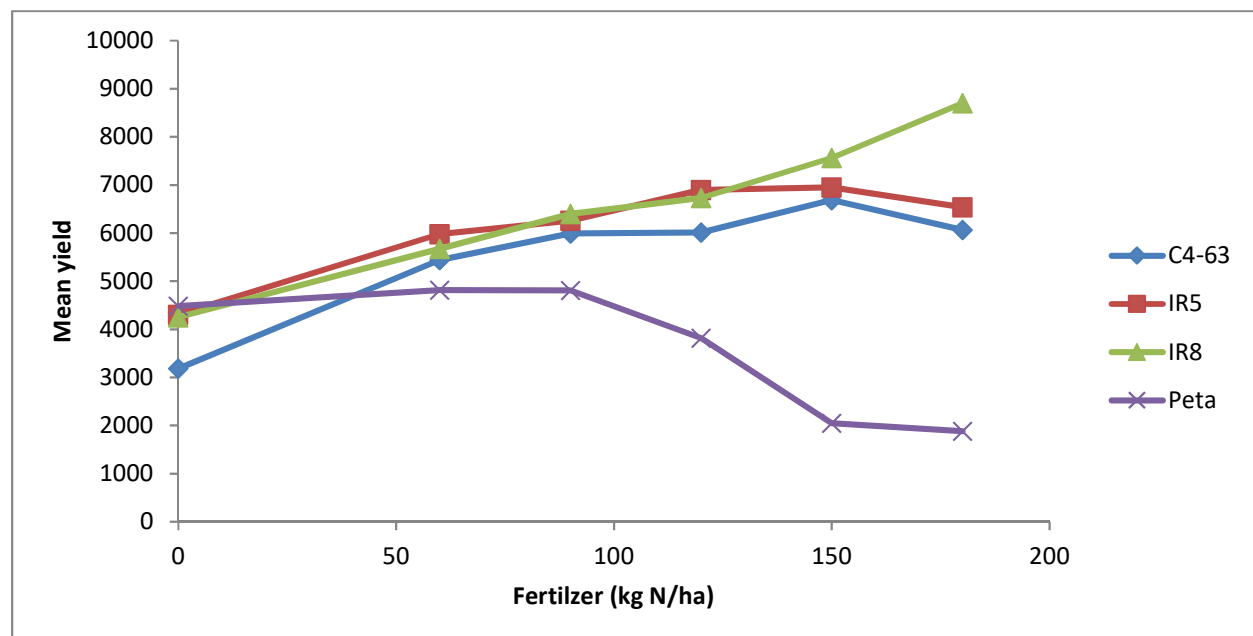
Table	Nitrogen	Variety	Nitrogen Variety
rep.	12	18	3
l.s.d.	342.7	399.7	898.1
d.f.	10	36	43.53
Except when comparing means with the same level(s) of Nitrogen			979.1
d.f.			36

Estimated stratum variances

Stratum	variance	effective d.f.	variance component
Rep	541288.3	2.000	16638.4
Rep.Nitrogen	141967.9	10.000	-51903.0
Rep.Nitrogen.Variety	349579.8	36.000	349579.8

Clearly the interpretation comes down to the fact that the response of rice to increasing amounts of nitrogen differs among the varieties ($F=13.22$, $P<0.001$).

A plot of means makes this very clear. The optimal amount of N to use clearly changes with the variety planted, with IR8 no worse than the others for all varieties.



The stratum variances illustrate an important feature of split-plot experiments. Intuitively, and mathematically, we would expect that the **block m.s.** would be larger than the **whole-plot**

Residual m.s., and that the **whole-plot Residual m.s.** would be larger than the **split-plot**

Residual m.s.. In this experiment, the **whole-plot Residual m.s.** () is *smaller* than the **split-plot**

Residual m.s.. Split-plots are smaller than whole-plots, so you would expect that variances to be based on smaller areas would be smaller than those based on larger areas. The three values are 541,288, **141,968** and **349,580**. This might necessitate a check on the experiment's protocols.

Before considering the next design, we look at the standard errors of differences carefully.

Standard errors of differences of means

Table	Nitrogen	Variety	Nitrogen Variety
rep.	12	18	3
s.e.d.	153.8	197.1	445.5
d.f.	10	36	43.53
Except when comparing means with the same level(s) of			
Nitrogen			482.8
d.f.			36

The column below Nitrogen is used for assessing differences in the overall nitrogen means:

0	60	90	120	150	180
4055.8	5478.2	5866.3	5864.4	5812.0	5796.8

There are 10 df for such comparisons. So for example, to compare the means for 60 versus 90 kg N/ha, a t test would be used, with

$$t = \frac{5866.3 - 5478.2}{153.8} = 2.52$$

and with df = 10, P = 0.030. A similar calculation could be made for varieties, using an s.e.d. value of 197.1 based on 36 df.

To compare say Peta with IR8, each with 90 kg N/ha, we are making a comparison at the same level of nitrogen, for which s.e.d. = 482.8 with 36 df. For this comparison

$$t = \frac{6400 - 4812}{153.8} = 3.29$$

and with df = 36, P = 0.002.

To compare say Peta with 90 kg N/ha to Peta with 120 kg N/ha, we are making a comparison at a different level of nitrogen, for which s.e.d. = 445.5 with 43.53 df. This is only an approximate t value, for which

$$t = \frac{4812 - 3816}{445.5} = 2.23$$

and with df = 43.53, P = 0.030.

Alternatively, print out l.s.d. values and use them in a similar way.

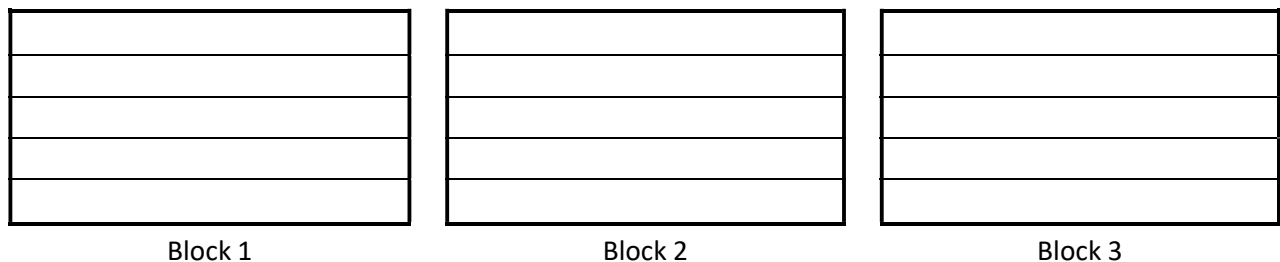
Strip-plot (or split-block) design

The next design is useful when there are two or more treatment factors and the total number of treatment combinations is too large to achieve a simple RCB layout. The difference between a split-plot and strip-plot design is that, with the latter design, the precision attached to the interaction takes precedence over that for both main effects. Recall that in the split-plot analysis the df for the split-plot Residual is the same as for the interaction.

The design is a three-stage one but results in 4 strata.

1. Form **blocks** in the field; these form the first strata.
2. In each block, form (say) *horizontal* strips to which one of the treatment factors, A (say) with a levels, is randomized into;
3. Again, in each block, form *vertical* strips to which the second treatment, B (say) with b levels, is applied to.

Pictorially, randomize the levels of A into these strips:

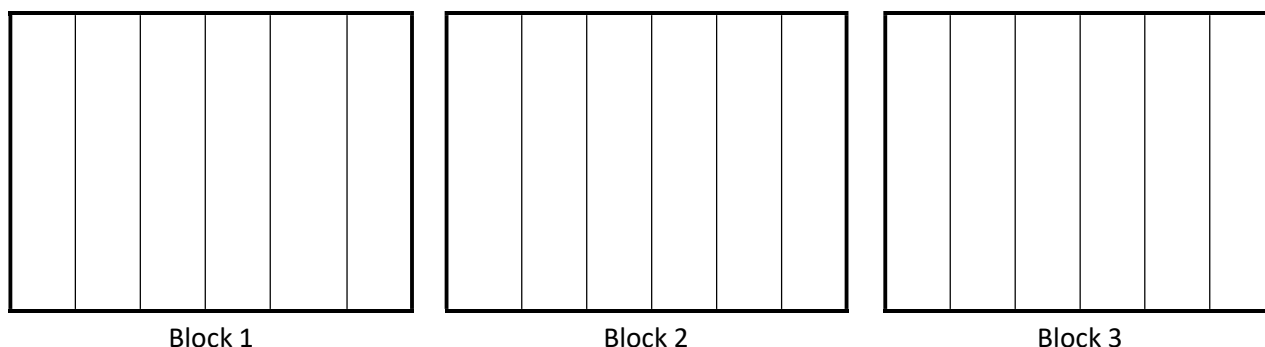


There is one replicate of each A treatment in each of r replicate blocks, hence horizontal strips form the replicates that allow A to be tested via an RCB ANOVA structure:

Component	df
Block stratum	
Block	$(r-1)$
Block.A stratum	
A	$(a-1)$
Residual (1) = Block.A	$(r-1)(a-1)$

In GenStat, this part of the Block Structure is **Block/A = Block + Block.A**.

Next, randomize the levels of B into these strips:



There is one replicate of each B treatment in each of r replicate blocks, hence vertical strips form the replicates that allow B to be tested via an RCB ANOVA structure:

Component	df
Block stratum	
Block	$(r-1)$
Block.B stratum	
B	$(b-1)$
Residual (2) = Block.B	$(r-1)(b-1)$

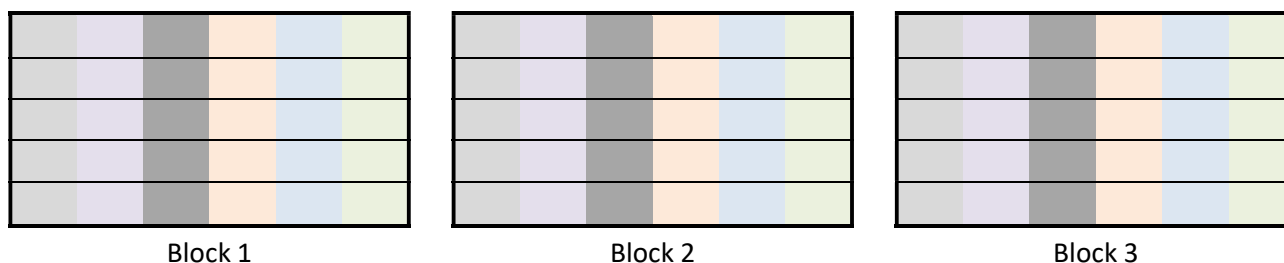
In GenStat, this part of the Block Structure is **Block/B = Block + Block.B**. There is no hierarchy between the two stages of randomization, hence both strata have equal priority in the analysis.

So to date we have discussed the formation of 3 strata in the analysis, and these are obtained in GenStat by combining the two concepts. When a factor occurs twice, the second occurrence is simply ignored:

$$\text{Block/A} + \text{Block/B} = (\text{Block} + \text{Block.A}) + (\text{Block} + \text{Block.B}) = \text{Block} + \text{Block.A} + \text{Block.B}$$

So to the final stratum.

In each block the two stripping processes give rise to the final field layout. Vertical strips have been colored to make identification easier.



The intersection of a vertical and horizontal strip is a single plot that received one level of factor A and one of factor B. There will be one replicate of each AB combination, one from each block. So the 3-stage randomization has induced a fourth stratum. The structure of this stratum could be omitted in GenStat (which always allows the final stratum to be omitted, adding it in if it is). But we simply provide a term that indexes over all blocks, levels of A and levels of B, so

Block.A.B. GenStat's Block Structure is, in full, **Block + Block.A + Block.B + Block.A.B.**

However, we have seen that $A + B + A.B$ can be simplified to $A*B$. So using GenStat's rules, the Block Structure can be simplified to **Block/(A*B)**

which intuitively makes sense: firstly form blocks, then within each block form strips into which factor A is randomized, at the same time form strips into which factor B is randomized, in such a way as they intersect – hence the $(A*B)$ part of the Block Structure, with neither factor A nor factor B taking priority.

So the structure of the ANOVA is the following. The order of the Block.A and Block.B strata depends only on the order the two factors are entered in the model; one is not split within the other, they are both stripped within blocks and both are at equal levels in the analysis.

Component	df
Block stratum	
Block	$(r-1)$
Block.A stratum	
A	$(a-1)$
Residual (1) = Block.A	$(r-1)(a-1)$
Block.B stratum	
B	$(b-1)$
Residual (2) = Block.B	$(r-1)(b-1)$
Block.A.B stratum	
A.B	$(b-1)$
Residual (3) = Block.A.B	$(r-1)(a-1)(b-1)$

Grain yield (kg/ha) of rice with 6 varieties and 3 levels of nitrogen, set out in 3 blocks as a strip plot, Page 110 of G&G

N (kg/ha)	Block 1	Block 2 IR8	Block 3
0	2373	3958	4384
60	4076	6431	4889
120	7254	6808	8582
	IR127-80		
0	4007	5795	5001
60	5630	7334	7177
120	7053	8284	6297
	IR305-4-12		
0	2620	4508	5621
60	4676	6672	7019
120	7666	7328	8611
	IR400-2-5		
0	2726	5630	3821
60	4838	7007	4816
120	6881	7735	6667
	IR665-58		
0	4447	3276	4582
60	5549	5340	6011
120	6880	5080	6076
	Peta		
0	2572	3724	3326
60	3896	2822	4425
120	1556	2706	3214

Analysis of Variance

Available Data: Nitrogen
Rep
Variety

Design: General Analysis of Variance

Y-Variate: Yield

Treatment Structure: Nitrogen*Variety

Block Structure: Rep/(Nitrogen*Variety)

Analysis of variance

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Rep stratum	2	9220962.	4610481.		
Rep.Nitrogen stratum					
Nitrogen	2	50676061.	25338031.	34.07	0.003
Residual	4	2974908.	743727.	1.81	
Rep.Variety stratum					
Variety	5	57100201.	11420040.	7.65	0.003
Residual	10	14922619.	1492262.	3.63	
Rep.Nitrogen.Variety stratum					
Nitrogen.Variety	10	23877979.	2387798.	5.80	<.001
Residual	20	8232917.	411646.		
Total	53	167005649.			

Message: the following units have large residuals.

Rep 1 Nitrogen 120 Variety Peta	-901.	s.e.	390.
Rep 2 Nitrogen 60 Variety IR8	818.	s.e.	390.
Rep 2 Nitrogen 60 Variety Peta	-1005.	s.e.	390.

Tables of means

Variate: Yield

Grand mean 5290.

Nitrogen	0	60	120			
	4021.	5478.	6371.			
Variety	IR127-80	IR305-4-12	IR400-2-5	IR665-58	IR8	Peta
	6286.	6080.	5569.	5249.	5417.	3138.

Nitrogen	Variety	IR127-80	IR305-4-12	IR400-2-5	IR665-58	IR8
0		4934.	4250.	4059.	4102.	3572.
60		6714.	6122.	5554.	5633.	5132.
120		7211.	7868.	7094.	6012.	7548.
Nitrogen	Variety	Peta				
0		3207.				
60		3714.				
120		2492.				

Standard errors of differences of means

Table	Nitrogen	Variety	Nitrogen Variety
rep.	18	9	3
s.e.d.	287.5	575.9	742.6
d.f.	4	10	22.29
Except when comparing means with the same level(s) of			
Nitrogen			717.3
d.f.			20.90
Variety			558.0
d.f.			22.43

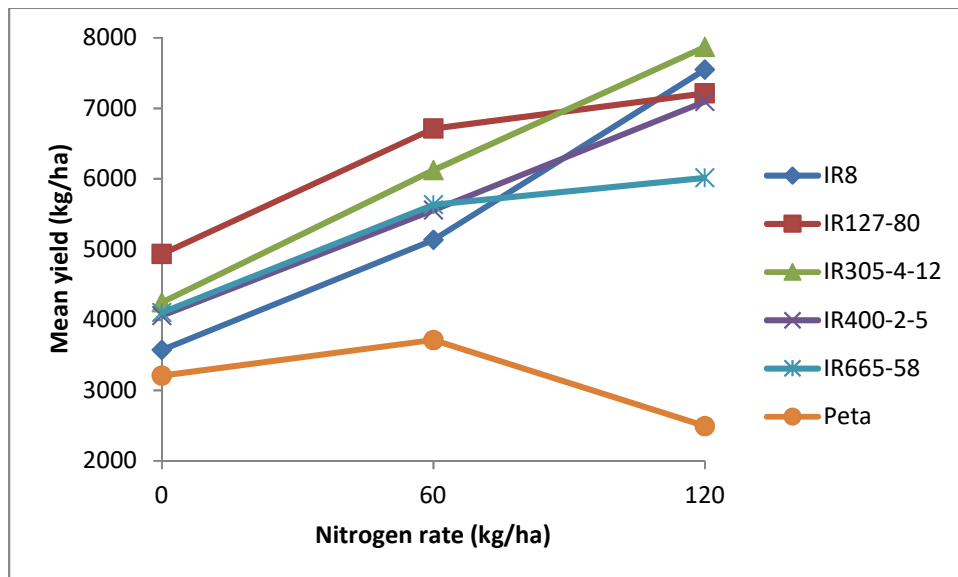
Least significant differences of means (5% level)

Table	Nitrogen	Variety	Nitrogen Variety
rep.	18	9	3
l.s.d.	798.1	1283.1	1538.9
d.f.	4	10	22.29
Except when comparing means with the same level(s) of			
Nitrogen			1492.2
d.f.			20.90
Variety			1155.9
d.f.			22.43

Estimated stratum variances

Stratum	variance	effective d.f.	variance component
Rep	4610481.2	2.000	154785.5
Rep.Nitrogen	743727.0	4.000	55346.9
Rep.Variety	1492261.9	10.000	360205.4
Rep.Nitrogen.Variety	411645.9	20.000	411645.9

Once again the interaction between varieties and nitrogen rates is highly significant ($F=5.80$, $P<0.001$). As with the previous data it is variety Peta whose response to N is different, as the plot of two-way means shows:



Notice also that none of the pairwise comparisons of these means is strictly a t test:

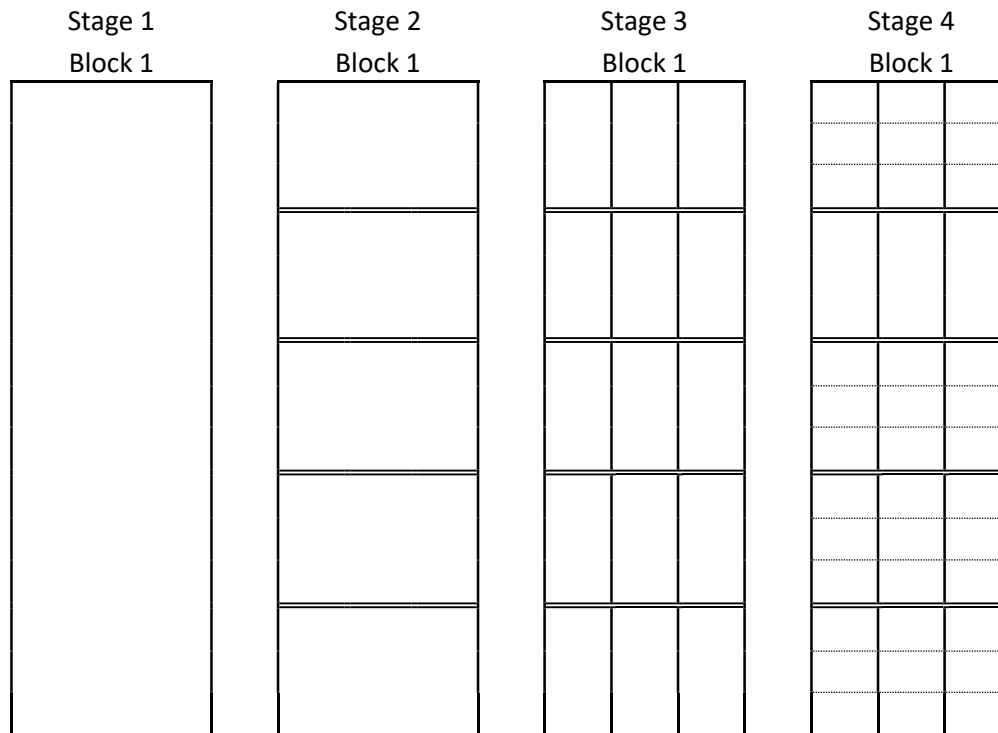
- 📊 To compare two variety means at a particular rate of N uses an s.e.d. value of 717.2 and the approximate t test would be based on 20.90 df;
- 📊 To compare two nitrogen means for a particular variety uses an s.e.d. value of 558.0 and the approximate t test would be based on 22.43 df;
- 📊 To compare two variety means, each at a different rate of N, uses an s.e.d. value of 742.6 and the approximate t test would be based on 22.29 df (though such a comparison would be unusual in practice).

Split-split-plot design in an RCB layout

The next design is useful when there are three or more treatment factors and too many combinations of two factors to form simple split plots. Generally, the factor applied to whole-plots is the least important treatment, or a treatment with levels whose means are quite different. The field layout comes about as a four stage randomization.

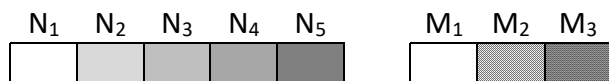
-  At stage 1, identify and construct an appropriate number of blocks;
-  At stage 2, form an appropriate number of large whole-plots in each block, to which one of the three treatment factors (the whole-plot treatment) is to be applied to;
-  At stage 3, form an appropriate number of smaller split-plots in each whole-plot, to which a second treatment factor (the split-plot treatment) is to be applied to;
-  At stage 4, form an appropriate number of even smaller split-split-plots in each split-plot, to which the last of the three treatment factors (the split-split-plot treatment) is to be applied to.

A feasible diagrammatic representation of this *for one of the blocks* is this:

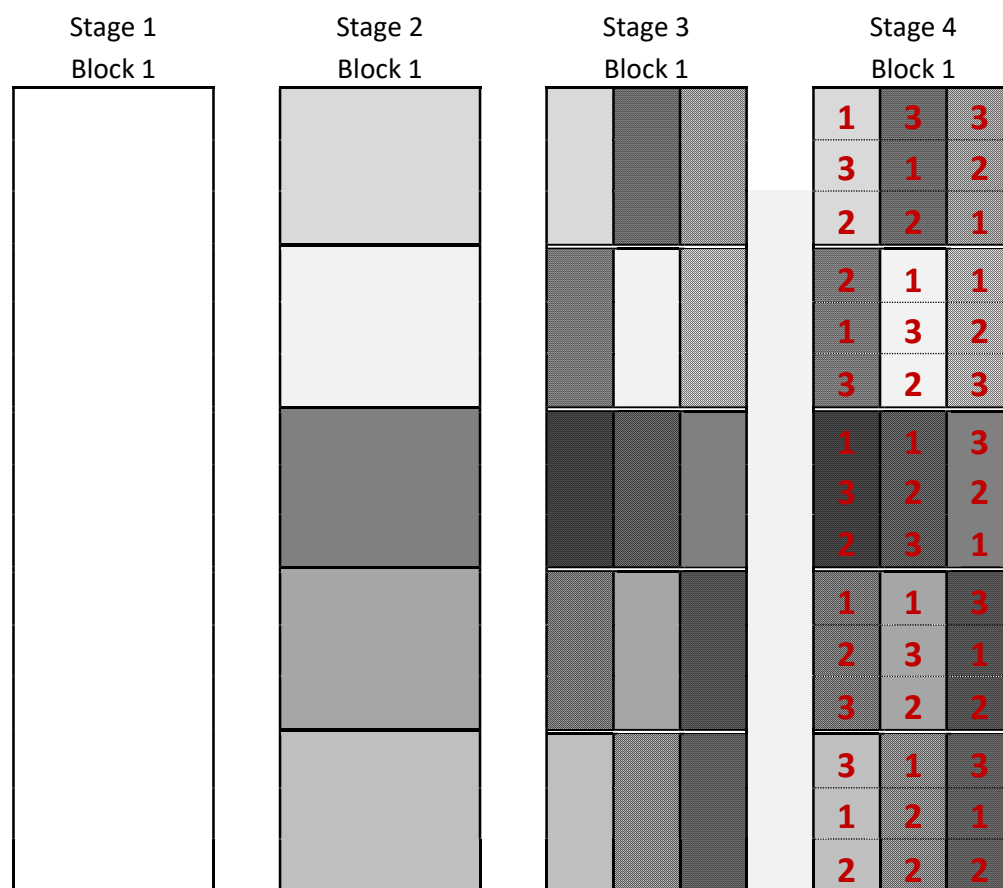


There are four strata in such a design. As we've seen, the Block Stratum contains just the Block component: there are no replicates of each block so there is no error term for constructing an F test for blocks. The other three strata do produce an error term, so we'll have three Residual components to define, each producing a different estimate of the variance of the yield from a plot in each stratum. But remember that blocks are generally regarded as random, so the **Block m.s.** does produce a Block Stratum Variance which can be printed as an option in the ANOVA.

Back to the randomization. Page 140 of G&G describes this nicely using an example with 3 replicate blocks, 5 rates of a nitrogen fertilizer as the whole-plot treatment, 3 management practices as the split treatment, and 3 varieties as the split-split treatment. Using the following shading patterns:



with red numbers for the varieties, the stages of randomization in their first replicate block are:



So what has been done in the field?

1. Firstly form blocks, so the Block Structure in GenStat starts with **Block**.
2. Next form whole-plots within each block, so **Block/whole-plot**.
3. Next form split-plots within each whole-plot, so **Block/whole-plot/split-plot**.
4. Finally form split-split-plots within each split-plot, so the final short-cut form of the Block Structure is **Block/whole-plot/split-plot/split-split-plot**.

Since the factor Nitrogen indexes over the number of whole-plots, we don't need a separate factor to identify the whole-plots, so in GenStat's Block Structure we can replace the factor **whole-plot** with **Nitrogen**.

Similarly, the factor **Management** can be used instead of a separate factor **split-plot**, and **Variety** can be used instead of a separate factor **split-split-plot**. So for a split-split plot analysis of variance an alternative Block Structure can use the actual treatment factors applied to each level: **Block/Nitrogen/Management/Variety**, or, since the final stratum can always be dropped, simply **Block/Nitrogen/Management**.

Using the full structure and GenStat's rule that A/B is a shortcut for A+A.B, we can expand the full Block Structure step by step and removing the repeated factor Block:

$$\begin{aligned}\text{Block/Nitrogen/(Management/Variety)} &= \text{Block/Nitrogen/(Management+ Management.Variety)} \\ &= \text{Block/Nitrogen+ Block/Nitrogen.(Management+ Management.Variety)} \\ &= (\text{Block+ Block.Nitrogen})+(\text{Block+ Block.Nitrogen.(Management+ Management.Variety)}) \\ &= \text{Block+ Block.Nitrogen+Block.Nitrogen.Management+ Block.Nitrogen.Management.Variety)}\end{aligned}$$

These are the four strata in the ANOVA, each with a different variance. The variety means are means of the individual smallest sized plots (the split-split plots) in the field, and hence the residual from the Block.Nitrogen.Management.Variety stratum is used to test the varieties. But the same is true for any mean involving varieties, whether they be two-way means (nitrogen × variety and management × variety means) or the three-way, nitrogen × management × variety means. These components all are tested in the lowest level stratum.

The management means, as well as the nitrogen $\times$ management means, are mean yields that had one management practice (or, in the case of the two-way means one management practice at one nitrogen rate applied). These are all based on split-split plots from Stage 3 of the field management, and both M and N.M components are tested using the residual from that stratum. So you can see that

1. the whole-plot treatment *Nitrogen* is tested against the residual from the *Block.Nitrogen* stratum,
2. the split-plot treatment *Management*, as well as the interaction of *Nitrogen* and *Management*, are tested against the residual from the *Block.Nitrogen.Management* stratum, and
3. the split-split-plot treatment *Variety*, as well as interactions involving *Variety* (so *Nitrogen.Variety*, *Management.Variety* and *Management.Nitrogen.Variety*), are tested against the residual from the *Block.Nitrogen.Management.Variety* stratum.

Using obvious notation for the numbers of levels ($r=3$ replicate blocks, $n=5$ levels of N, $m=3$ levels of management M and $v=3$ varieties V) the split-split-plot analysis.

Component	df
Block stratum	
Block	$(r-1)=2$
Block.Nitrogen stratum	
N	$(n-1) = 4$
Residual (1) = Block.N	$(r-1)(n-1) = 8$
Block.Nitrogen.Management stratum	
M	$(m-1) = 2$
N.M	$(n-1)(m-1) = 8$
Residual (2)	$n(r-1)(m-1) = 20$
Block.Nitrogen.Management.Variety stratum	
V	$(v-1) = 2$
N.V	$(n-1)(v-1) = 8$
M.V	$(m-1)(v-1) = 4$
N.M.V	$(n-1)(m-1)(v-1) = 16$
Residual (3)	$nm(r-1)(v-1) = 60$

Grain yields (t/ha) of 3 rice varieties under 3 management practices and 5 levels of nitrogen, using a split-split plot RCB, from G&G Page 143; nitrogen is whole-plot treatment, management the split-plot treatment and variety the split-split treatment

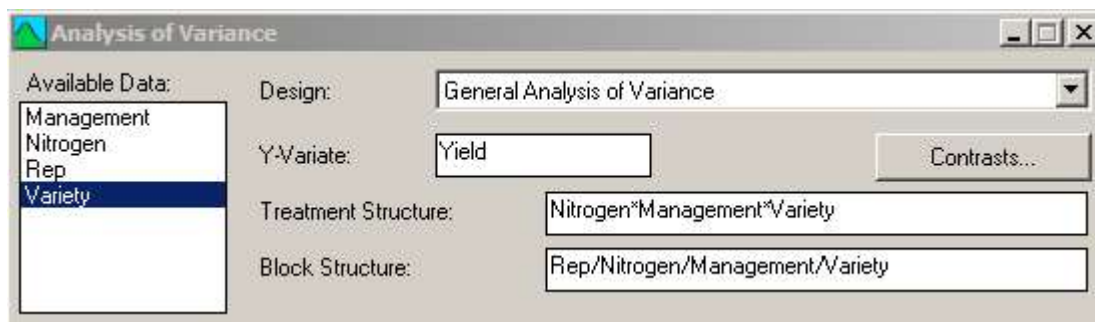
Management	Variety 1 replicate			Variety 2 replicate			Variety 3 replicate		
	1	2	3	1	2	3	1	2	3
0 kg N/ha									
Minimum	3.320	3.864	4.507	6.101	5.122	4.815	5.355	5.536	5.244
Optimum	3.766	4.311	4.875	5.096	4.873	4.166	7.442	6.462	5.584
Intensive	4.660	5.915	5.400	6.573	5.495	4.225	7.018	8.020	7.642
50 kg N/ha									
Minimum	3.188	4.752	4.756	5.595	6.780	5.390	6.706	6.546	7.092
Optimum	3.625	4.809	5.295	6.357	5.925	5.163	8.592	7.646	7.212
Intensive	5.232	5.170	6.046	7.016	7.442	4.478	8.480	9.942	8.714
80 kg N/ha									
Minimum	5.468	5.788	4.422	5.442	5.988	6.509	8.452	6.698	8.650
Optimum	5.759	6.130	5.308	6.398	6.533	6.569	8.662	8.526	8.514
Intensive	6.215	7.106	6.318	6.953	6.914	7.991	9.112	9.140	9.320
110 kg N/ha									

Minimum	4.246	4.842	4.863	6.209	6.768	5.779	8.042	7.414	6.902
Optimum	5.255	5.742	5.345	6.992	7.856	6.164	9.080	9.016	7.778
Intensive	6.829	5.869	6.011	7.565	7.626	7.362	9.660	8.966	9.128

140 kg N/ha

Minimum	3.132	4.375	4.678	6.860	6.894	6.573	9.314	8.508	8.032
Optimum	5.389	4.315	5.896	6.857	6.974	7.422	9.224	9.680	9.294
Intensive	5.217	5.389	7.309	7.254	7.812	8.950	10.360	9.896	9.712

The ANOVA is obtained easily enough:



Analysis of variance

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Rep stratum	2	0.7320	0.3660	0.66	
Rep.Nitrogen stratum					
Nitrogen	4	61.6408	15.4102	27.70	<.001
Residual	8	4.4514	0.5564	2.13	
Rep.Nitrogen.Management stratum					
Management	2	42.9361	21.4681	82.00	<.001
Nitrogen.Management	8	1.1030	0.1379	0.53	0.823
Residual	20	5.2363	0.2618	0.53	
Rep.Nitrogen.Management.Variety stratum					
Variety	2	206.0132	103.0066	207.87	<.001
Nitrogen.Variety	8	14.1445	1.7681	3.57	0.002
Management.Variety	4	3.8518	0.9629	1.94	0.115
Nitrogen.Management.Variety	16	3.6992	0.2312	0.47	0.954
Residual	60	29.7325	0.4955		
Total	134	373.5407			

Message: the following units have large residuals.

Rep 3 Nitrogen 140	0.386	s.e. 0.182
Rep 1 Nitrogen 0 Management Intensive Variety V2	1.164	s.e. 0.469
Rep 3 Nitrogen 50 Management Intensive Variety V2	-1.300	s.e. 0.469

Tables of means

Variate: Yield

Grand mean 6.554

Nitrogen	0	50	80	110	140
	5.385	6.220	6.996	6.937	7.234
Management	Intensive	Minimum	Optimum		
	7.277	5.900	6.486		
Variety	V1	V2	V3		
	5.127	6.396	8.140		
Nitrogen	Management	Intensive	Minimum	Optimum	
0		6.105	4.874	5.175	
50		6.947	5.645	6.069	
80		7.674	6.380	6.933	
110		7.668	6.118	7.025	
140		7.989	6.485	7.228	
Nitrogen	Variety	V1	V2	V3	
0		4.513	5.163	6.478	
50		4.764	6.016	7.881	
80		5.835	6.589	8.564	
110		5.445	6.925	8.443	
140		5.078	7.288	9.336	
Management	Variety	V1	V2	V3	
Intensive		5.912	6.910	9.007	
Minimum		4.413	6.055	7.233	
Optimum		5.055	6.223	8.181	
Nitrogen	Management	Variety	V1	V2	V3
0	Intensive		5.325	5.431	7.560
	Minimum		3.897	5.346	5.378
	Optimum		4.317	4.712	6.496
50	Intensive		5.483	6.312	9.045
	Minimum		4.232	5.922	6.781
	Optimum		4.576	5.815	7.817
80	Intensive		6.546	7.286	9.191
	Minimum		5.226	5.980	7.933
	Optimum		5.732	6.500	8.567
110	Intensive		6.236	7.518	9.251
	Minimum		4.650	6.252	7.453
	Optimum		5.447	7.004	8.625
140	Intensive		5.972	8.005	9.989

Minimum	4.062	6.776	8.618
Optimum	5.200	7.084	9.399

Standard errors of differences of means

Table	Nitrogen	Management	Variety	Nitrogen Management
rep.	27	45	45	9
s.e.d.	0.2030	0.1079	0.1484	0.2828
d.f.	8	20	60	22.26
Except when comparing means with the same level(s) of Nitrogen				0.2412
d.f.				20

Table	Nitrogen Variety	Management Variety	Nitrogen Management Variety
rep.	9	15	3
s.e.d.	0.3386	0.2360	0.5479
d.f.	43.48	79.29	82.25
Except when comparing means with the same level(s) of Nitrogen			0.5277
d.f.	60		79.29
Management		0.2570	
d.f.		60	
Nitrogen.Management			0.5748
d.f.			60
Nitrogen.Variety			0.5277
d.f.			79.29

Least significant differences of means (5% level)

Table	Nitrogen	Management	Variety	Nitrogen Management
rep.	27	45	45	9
l.s.d.	0.4682	0.2250	0.2969	0.5862
d.f.	8	20	60	22.26
Except when comparing means with the same level(s) of Nitrogen				0.5032
d.f.				20

Table	Nitrogen Variety	Management Variety	Nitrogen Management Variety
rep.	9	15	3
l.s.d.	0.6826	0.4697	1.0900
d.f.	43.48	79.29	82.25
Except when comparing means with the same level(s) of Nitrogen			1.0502
d.f.	60		79.29
Management		0.5142	
d.f.		60	

Nitrogen.Management		1.1497	
d.f.		60	
Nitrogen.Variety		1.0502	
d.f.		79.29	
Estimated stratum variances			
Variate: Yield			
Stratum	variance	effective d.f.	variance component
Rep	0.3660	2.000	-0.0042
Rep.Nitrogen	0.5564	8.000	0.0327
Rep.Nitrogen.Management	0.2618	20.000	-0.0779
Rep.Nitrogen.Management.Variety	0.4955	60.000	0.4955

Interpretation:

✚ The stratum variances indicate less variation in yield than the experimenters might have expected. The **Block m.s.** is smaller than the whole-plot **Residual m.s.** (0.3660 compared to 0.5564) and results in a negative estimate of variance – which simply implies there is no change in yield across blocks. Remember that blocks are generally used to control for variance, and the assumption is that plots in each block have some differing effect on growing conditions.

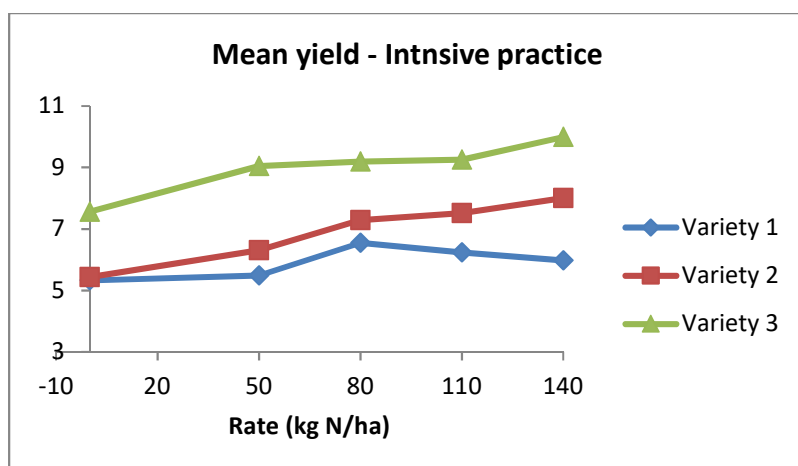
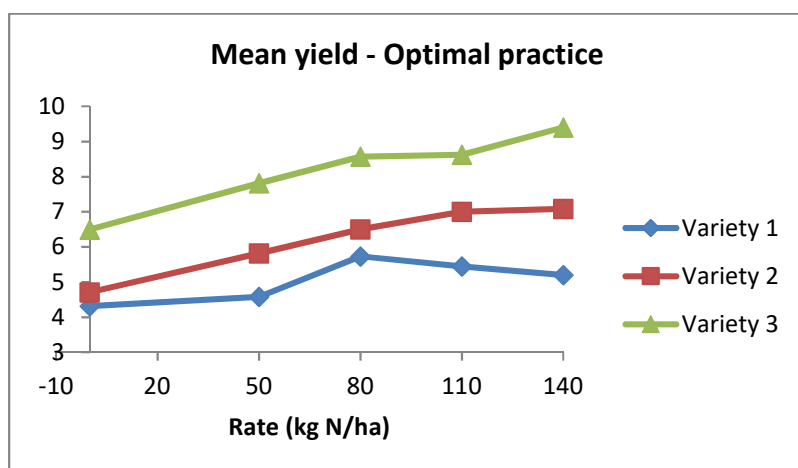
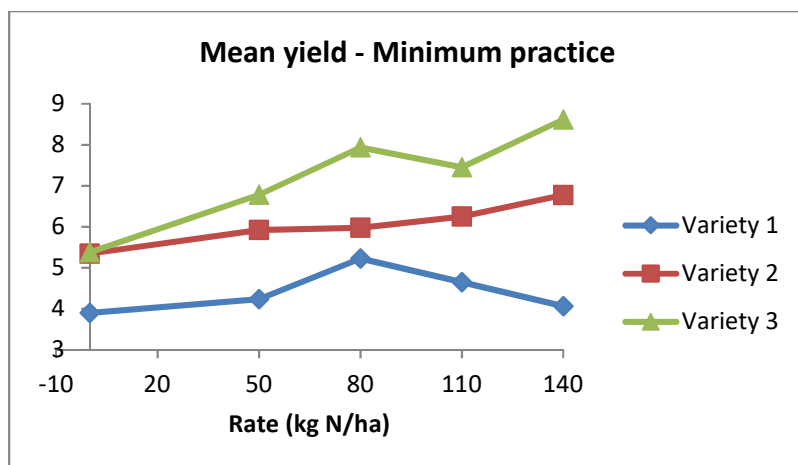
✚ Note also that the variance of the split-split plots is larger than that of the split-plots (0.4955 compared to 0.2618), again something that is not expected practically or mathematically.

✚ There is no 3-factor interaction ($F=0.47$, $P=0.954$). Of the two factor interactions, the effect on rice of nitrogen changes with the variety ($F=3.57$, $P=0.002$). The first of these

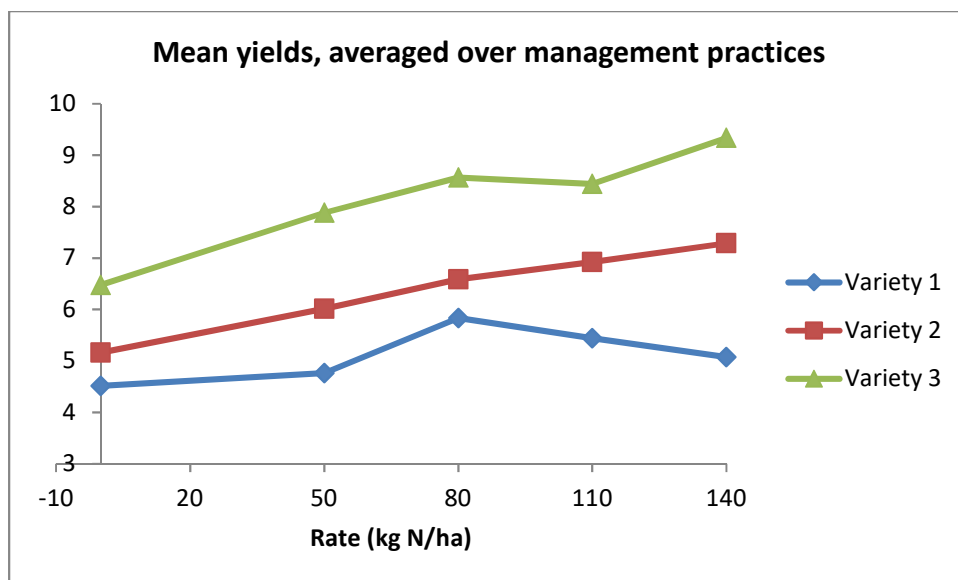
two statements implies that the interaction of nitrogen and variety is consistent across management practices.

📊 The management practices have a definite effect on rice yields ($F=82.0$, $P<0.001$), with intensive farming producing significantly better yields (mean yields are Management: 7.277, Optimum: 6.486, Minimum: 5.900 with an l.s.d. value of 0.227).

📊 The plot of three-way means illustrates these findings. The three-factor interaction compares the trends in N for the three varieties across the three management practices, so graphically all three plots should appear very similar (within statistical variation):



Since this interaction is not significant, we can present a single plot with means averaged across management practices:



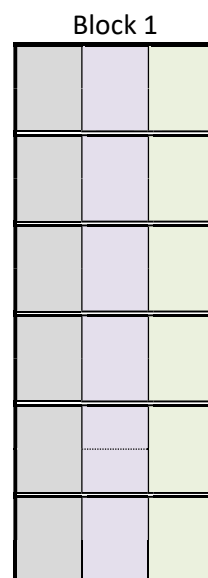
Remember also that comparisons between means are judged using l.s.d. values, or P values calculated by taking mean differences and dividing by the s.e.d. values. In these two tables in GenStat's output, only those comparisons whose degrees of freedom are integers are strictly t tests; the rest are approximations.

Strip-split-plot design

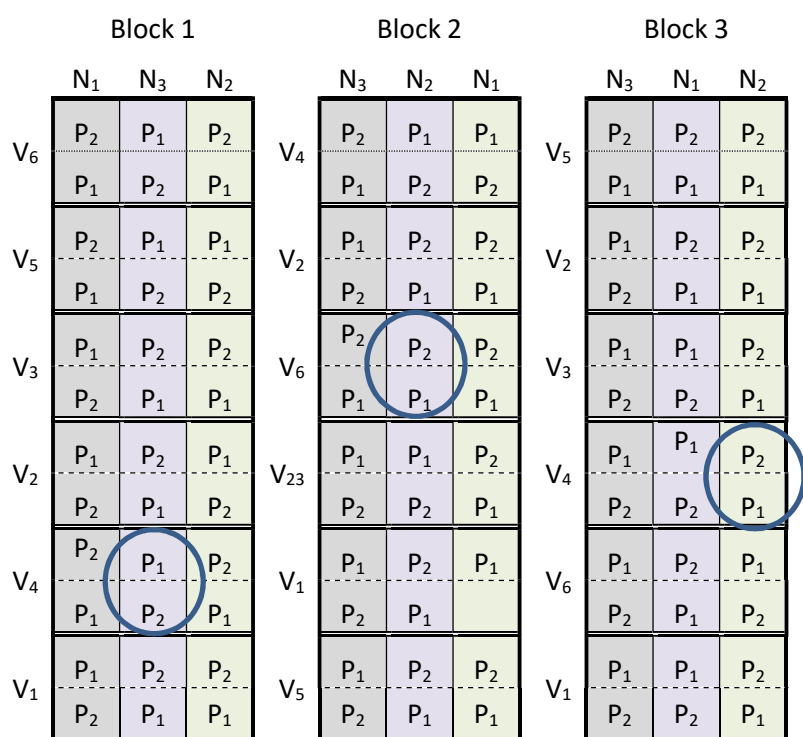
The next design is a refinement on the strip-plot design to allow for a third splitting of the small units resulting from that design. Firstly, recall that for the strip-plot design we have this layout with 4 strata (the strata being blocks, horizontal strips in each block, vertical strips in each block and the intersection of a vertical and horizontal strip in each block). Here is a typical block:

GenStat's Block Structure for the strip-plot design is:

$$\text{Block} + \text{Block.A} + \text{Block.B} + \text{Block.A.B} = \text{Block}/(\text{A} * \text{B})$$



The refinement to this design that gives rise to the new design allowing a third splitting of these small units into which a third treatment is randomized. The dataset on Page 155 of G&G has 6 varieties (V) applied to horizontal strips, 3 nitrogen rates (N) applied to vertical strips and 2 planting methods (V) **randomized into every intersection plot**, with three block replicates (3 such plots are circles to highlight the independent randomizations taking place in these intersection plots):



So how is GenStat's Block Structure modified? The fact is that the current structure

Block/(Variety*Nitrogen) defines the first 4 strata (blocks, horizontal plots to which varieties are randomized, vertical plots to which nitrogen rates are randomized and intersection plots which form the replicates for the interaction of varieties with nitrogen) and, using GenStat's rule that the last stratum can be omitted, **Block/(Variety*Nitrogen)** is sufficient to define the analysis.

Strictly speaking, though, each intersection plot is split into small units for the randomization of planting methods, and, using the same intuitive approach that GenStat employs, this implies that the full Block Structure with all 5 strata defined is:

Block/(Variety*Nitrogen)/Planting_method

So this philosophy intuitively leads to F tests in separate strata:

1. Blocks are unreplicated and appear (alone) in the first stratum.
2. Varieties are randomized into horizontal plots, one replicate of each variety in each block (just as in an RCB), so the main effect for Variety is tested in a stratum using the Rep.Variety residual to form the F test.
3. Similarly nitrogen rates are randomized into vertical plots (just as in an RCB), one replicate of each nitrogen rate in each block, so the main effect for Nitrogen is tested in a separate stratum using the Rep.Nitrogen residual to form the F test.

1.	Block stratum	
	Block	$(b-1)$
2.	Block.Variety stratum	
	Variety	$(v-1)$
	Residual = Block.Variety	$(b-1)(v-1)$
3.	Block.Nitrogen stratum	
	Nitrogen	$(n-1)$
	Residual = Block. Nitrogen	$(b-1)(n-1)$
4.	Block.Variety.Nitrogen stratum	
	Variety.Nitrogen	$(v-1)(n-1)$
	Residual	$(b-1)(v-1)(n-1)$
5.	Block.Variety.Nitrogen.Planting stratum	
	Planting	$(p-1)$
	Variety.Planting	$(v-1)(p-1)$
	Nitrogen.Planting	$(n-1)(p-1)$
	Variety.Nitrogen.Planting	$(v-1)(n-1)(p-1)$
	Residual	$vn(b-1)(p-1)$

4. The intersection plots had a particular variety and nitrogen rate applied, so these units form the replicates for testing the Variety.Nitrogen interaction, the Residual being the interaction Block.Variety.Nitrogen and is used to construct the F test for the Variety.Nitrogen interaction.
5. The main effect (Planting method) and any interaction involving planting method are based on split-plot replicates, and hence these all appear in a final split-plot unit stratum.

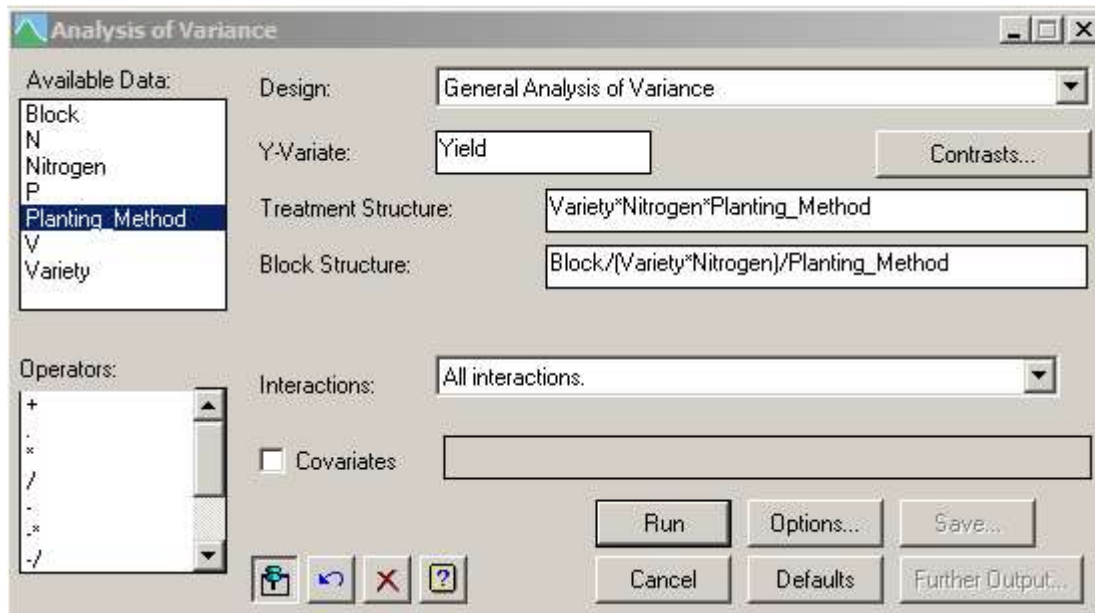
Dataset from G&G, Page 155, rearranged in field layout (assuming the randomized layout above

is the one actually used for the experiment). N represents nitrogen rates 0, 60, 120 kg N / ha;

P1 is broadcast, P2 transplanted; the 6 varieties are IR8, IR127-8-1-10, IR305-4-12-1-3, IR400-2-

5-3-3-2, IR665-58 and Peta:

	N ₁	N ₃	N ₂		N ₃	N ₂	N ₁		N ₃	N ₁	N ₂
V ₆	4535 P2	1556 P1	4627 P2	V ₄	9838 P2	7007 P1	5630 P1	V ₅	6564 P2	3739 P2	4666 P2
	2572 P1	5374 P2	3896 P1		7735 P1	6928 P2	6200 P2		6076 P1	4582 P1	6011 P1
V ₅	4655 P2	6880 P1	5549 P1	V ₂	8284 P1	7424 P2	4885 P2	V ₂	6297 P1	4583 P2	5377 P2
	4447 P1	6995 P2	4646 P2		7648 P2	7334 P1	5795 P1		5736 P2	5001 P1	7177 P1
V ₃	2620 P1	8632 P2	4946 P2	V ₆	7218 P2	4461 P2	5457 P2	V ₃	8611 P1	5621 P1	6142 P2
	4527 P2	7666 P1	4676 P1		2706 P1	2822 P1	3724 P1		7416 P2	3628 P2	7019 P1
V ₂	4007 P1	6440 P2	5630 P1	V ₃	7328 P1	6672 P1	4866 P2	V ₄	6667 P1	3821 P1	4829 P2
	4035 P2	7053 P1	3728 P2		7101 P2	7611 P2	4508 P1		7253 P2	4038 P2	4816 P1
V ₄	5274 P2	6881 P1	4878 P2	V ₁	6808 P1	7502 P2	3958 P1	V ₆	3214 P1	3537 P2	4425 P1
	2726 P1	6545 P2	4838 P1		6353 P2	6431 P1	3528 P2		6369 P2	3326 P1	4774 P2
V ₁	2373 P1	6661 P2	3085 P2	V ₅	5080 P1	5006 P2	2796 P2	V ₁	8582 P1	4384 P1	4362 P2
	2293 P2	7254 P1	4076 P1		4486 P2	5340 P1	3276 P1		7759 P2	2538 P2	4889 P1



Analysis of variance

Variate: Yield

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Block stratum	2	15289498.	7644749.		
Block.Variety stratum					
Variety	5	49119270.	9823854.	3.68	0.038
Residual	10	26721828.	2672183.	2.80	
Block.Nitrogen stratum					
Nitrogen	2	116489166.	58244583.	36.62	0.003
Residual	4	6361491.	1590373.	1.66	
Block.Variety.Nitrogen stratum					
Variety.Nitrogen	10	24595731.	2459573.	2.57	0.034
Residual	20	19106733.	955337.	2.27	
Block.Variety.Nitrogen.Planting_Method stratum					
Planting_Method	1	723079.	723079.	1.71	0.199
Variety.Planting_Method	5	23761441.	4752288.	11.27	<.001
<i>Nitrogen.Planting_Method</i>	2	2468132.	1234066.	2.93	0.066
Variety.Nitrogen.Planting_Method	10	7512072.	751207.	1.78	0.100
Residual	36	15179354.	421649.		
Total	107	307327795.			

Message: the following units have large residuals.

Block 2 Variety IR665-58	-1168.	s.e. 497.
Block 1 Variety Peta Nitrogen 120	-878.	s.e. 421.
Block 2 Variety IR8 Nitrogen 60	924.	s.e. 421.
Block 2 Variety IR8 Nitrogen 120	-882.	s.e. 421.
Block 2 Variety Peta Nitrogen 60	-1159.	s.e. 421.
Block 1 Variety IR305-4-12-1-13 Nitrogen 0 Planting_Method Broadcast	-908.	s.e. 375.
Block 1 Variety IR305-4-12-1-13 Nitrogen 0 Planting_Method Transplanted	908.	s.e. 375.
Block 3 Variety IR305-4-12-1-13 Nitrogen 0 Planting_Method Broadcast	1042.	s.e. 375.
Block 3 Variety IR305-4-12-1-13 Nitrogen 0 Planting_Method Transplanted	-1042.	s.e. 375.

Tables of means

Variate: Yield

Grand mean 5372.

Variety	IR8	IR127-8-1-10	IR305-4-12-1-13	IR400-2-5-3-3-2
	5158.	5913.	6088.	5884.
Variety	IR665-58	Peta		
	5044.	4144.		
Nitrogen	0	60	120	
	4097.	5378.	6641.	
Planting_Method	Broadcast	Transplanted		
	5290.	5454.		
Variety	Nitrogen	0	60	120
IR8		3179.	5058.	7236.
IR127-8-1-10		4718.	6112.	6910.
IR305-4-12-1-13		4295.	6178.	7792.
IR400-2-5-3-3-2		4615.	5549.	7486.
IR665-58		3916.	5203.	6013.
Peta		3859.	4168.	4406.
Variety	Planting_Method	Broadcast	Transplanted	
IR8		5417.	4898.	
IR127-8-1-10		6286.	5540.	
IR305-4-12-1-13		6080.	6097.	
IR400-2-5-3-3-2		5569.	6198.	
IR665-58		5249.	4839.	
Peta		3138.	5150.	
Nitrogen	Planting_Method	Broadcast	Transplanted	
0		4021.	4173.	
60		5478.	5277.	
120		6371.	6910.	

Variety	Nitrogen	Planting_Method	Broadcast	Transplanted
IR8	0		3572.	2786.
	60		5132.	4983.
	120		7548.	6924.
IR127-8-1-10	0		4934.	4501.
	60		6714.	5510.
	120		7211.	6608.
IR305-4-12-1-13	0		4250.	4340.
	60		6122.	6233.
	120		7868.	7716.
IR400-2-5-3-3-2	0		4059.	5171.
	60		5554.	5545.
	120		7094.	7879.
IR665-58	0		4102.	3730.
	60		5633.	4773.
	120		6012.	6015.
Peta	0		3207.	4510.
	60		3714.	4621.
	120		2492.	6320.

Standard errors of means

Table	Variety	Nitrogen	Planting_Method	Variety Nitrogen
rep.	18	36	54	6
e.s.e.	385.3	210.2	88.4	521.8
d.f.	10	4	36	24.63
Except when comparing means with the same level(s) of				
Variety				420.6
d.f.				22.86
Nitrogen				504.6
d.f.				23.42

Table	Variety Planting_Method	Nitrogen Planting_Method	Variety Nitrogen Planting_Method
rep.	9	18	3
e.s.e.	414.6	236.4	585.3
d.f.	13.31	6.35	37.29
Except when comparing means with the same level(s) of			
Variety	216.4		497.1
d.f.	36		40.57
Nitrogen		153.1	570.0
d.f.		36	36.34
Variety.Nitrogen			374.9
d.f.			36
Variety.Planting_Method			497.1
d.f.			40.57
Nitrogen.Planting_Method			570.0
d.f.			36.34

Standard errors of differences of means

Table	Variety	Nitrogen	Planting_Method	Variety Nitrogen
rep.	18	36	54	6
s.e.d.	544.9	297.2	125.0	737.9
d.f.	10	4	36	24.63
Except when comparing means with the same level(s) of				
Variety				594.7
d.f.				22.86
Nitrogen				713.6
d.f.				23.42

Table	Variety Planting_Method	Nitrogen Planting_Method	Variety Nitrogen Planting_Method
rep.	9	18	3
s.e.d.	586.3	334.3	827.7
d.f.	13.31	6.35	37.29
Except when comparing means with the same level(s) of			
Variety	306.1		703.0
d.f.	36		40.57
Nitrogen		216.4	806.1
d.f.		36	36.34
Variety.Nitrogen			530.2
d.f.			36
Variety.Planting_Method			703.0
d.f.			40.57
Nitrogen.Planting_Method			806.1
d.f.			36.34

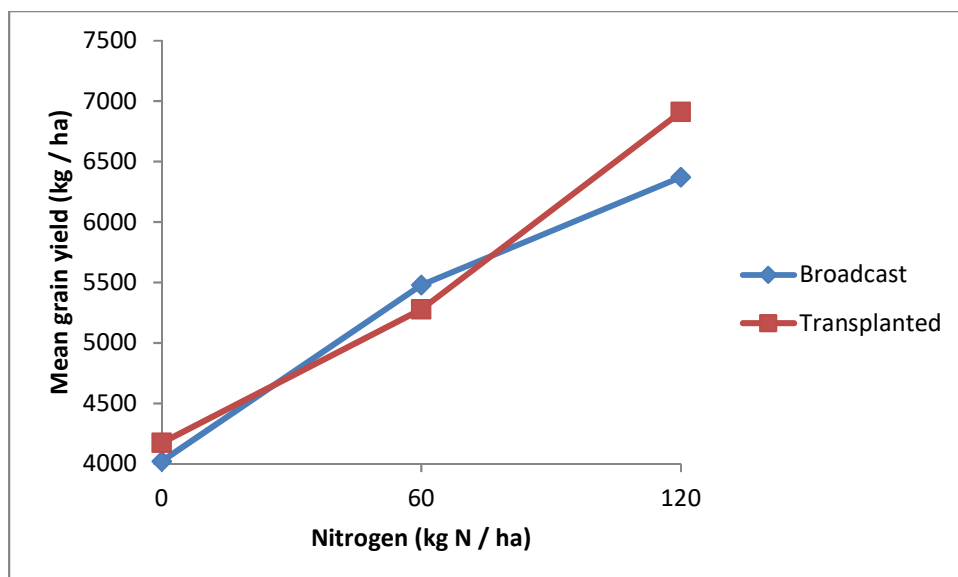
Estimated stratum variances

Variate: Yield

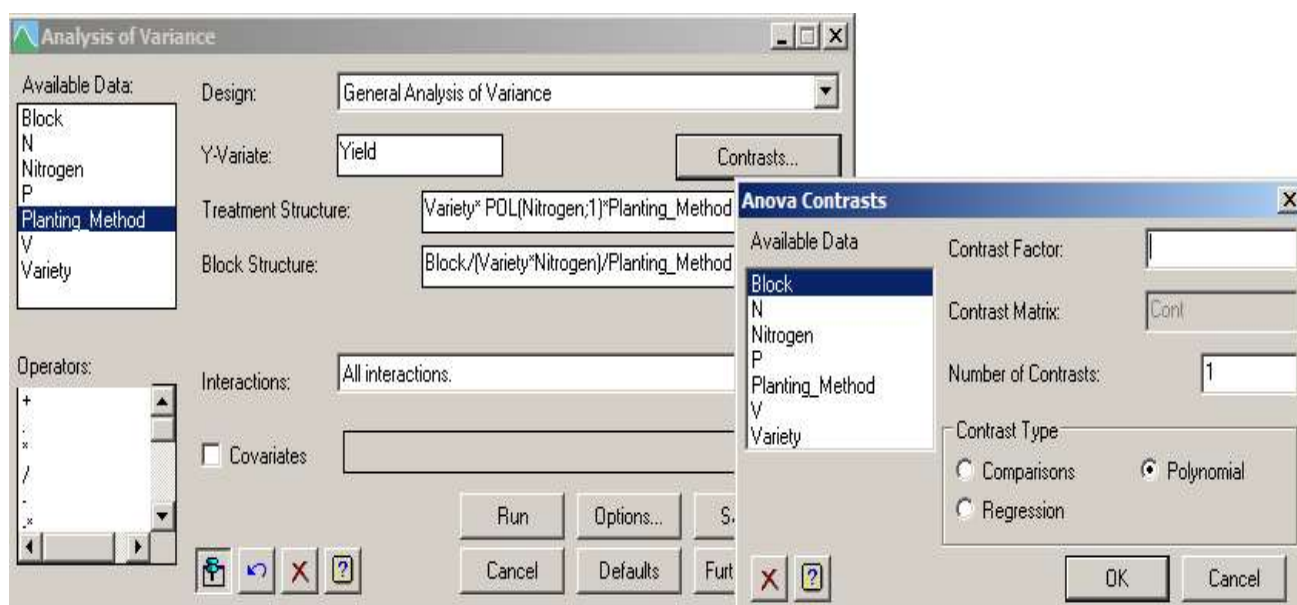
Stratum	variance	effective d.f.	variance component
Block	7644749.1	2.000	120486.9
Block.Variety	2672182.8	10.000	286141.0
Block.Nitrogen	1590372.8	4.000	52919.7
Block.Variety.Nitrogen	955336.7	20.000	266844.0
Block.Variety.Nitrogen.Planting_Method	421648.7	36.000	421648.7

Interpretation:

1. The large residuals flagged at the end of the ANOVA are informative. The yield in block 3 that grew IR305-4-12-1-13 (V2) with no nitrogen is extraordinarily large (5621 kg N / ha); the mean for the two planting methods with this combination is only 4,250-4,340 kg N / ha. The standardized residual for these two plots is $\pm 1042/375 = 2.78$; only 0.5% of all residuals should be this extreme. However, we have no way of checking the data, so take the analysis on face value.
2. The significant main effects for variety and nitrogen are irrelevant given the significant two-way interactions (and no significant three-way interaction). In fact these varieties and nitrogen rates were chosen presumably in the knowledge that there were differences, and that may be (partly) why they were chosen as the two whole-plot treatments to apply to horizontal and vertical strips. Attention should focus on whether these expected differences were consistent across the levels of the other factor.
3. The nitrogen effect is different across varieties ($P=0.034$) and almost significant across planting methods ($P=0.066$). A plot of means is as follows:



Clearly there is a strong linear trend for both planting methods, a feature that could have been extracted as part of the ANOVA:

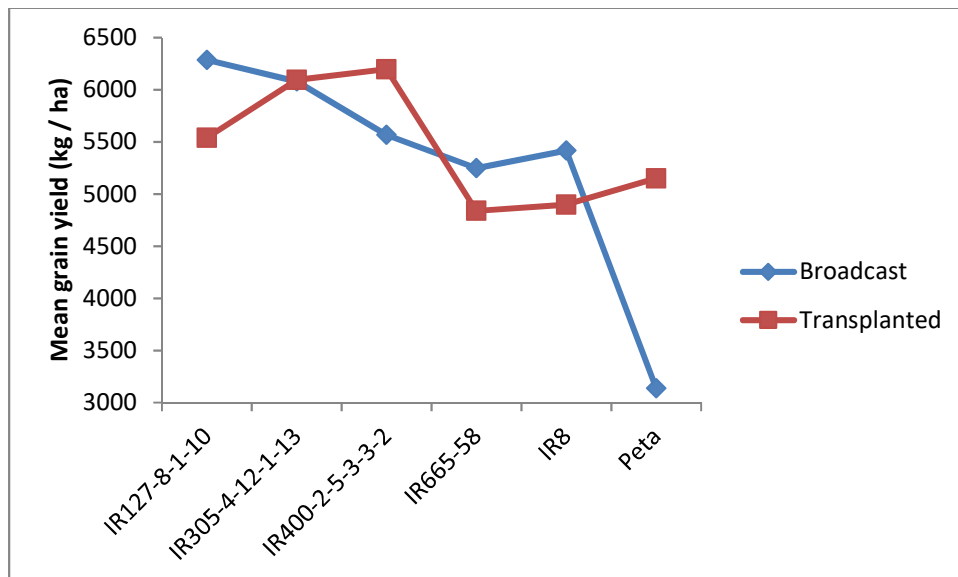


Analysis of variance

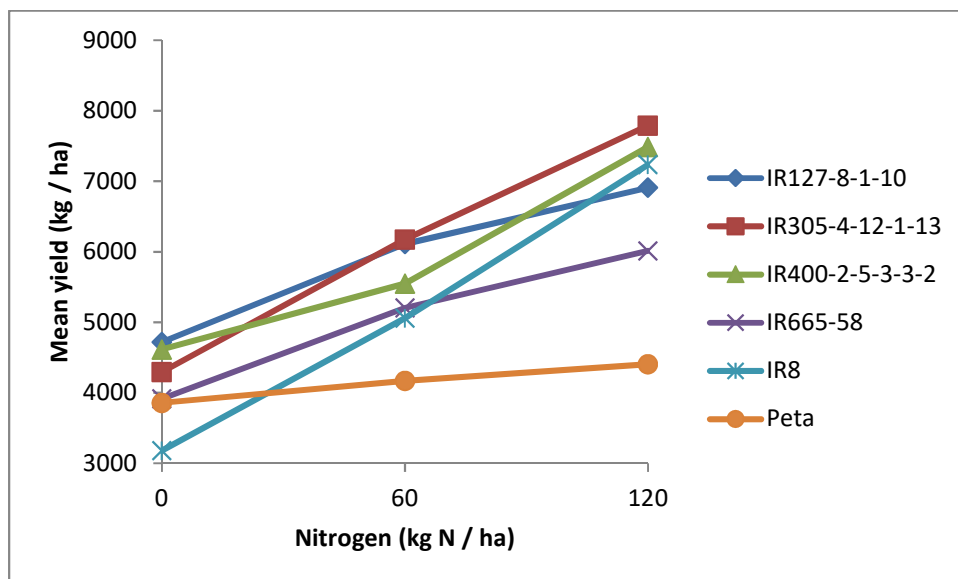
Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Block stratum	2	15289498.	7644749.		
Block.Variety stratum					
Variety	5	49119270.	9823854.	3.68	0.038
Residual	10	26721828.	2672183.	2.80	
Block.Nitrogen stratum					
Nitrogen	2	116489166.	58244583.	36.62	0.003
Lin	1	116487216.	116487216.	73.25	0.001
Deviations	1	1950.	1950.	0.00	0.974
Residual	4	6361491.	1590373.	1.66	
Block.Variety.Nitrogen stratum					
Variety.Nitrogen	10	24595731.	2459573.	2.57	0.034
Variety.Lin	5	22843042.	4568608.	4.78	0.005
Deviations	5	1752688.	350538.	0.37	0.865
Residual	20	19106733.	955337.	2.27	
Block.Variety.Nitrogen.Planting_Method stratum					
Planting_Method	1	723079.	723079.	1.71	0.199
Variety.Planting_Method	5	23761441.	4752288.	11.27	<.001
Nitrogen.Planting_Method	2	2468132.	1234066.	2.93	0.066
Lin.Planting_Method	1	674154.	674154.	1.60	0.214
Deviations	1	1793978.	1793978.	4.25	0.046
Variety.Nitrogen.Planting_Method	10	7512072.	751207.	1.78	0.100
Variety.Lin.Planting_Method	5	4382437.	876487.	2.08	0.091
Deviations	5	3129635.	625927.	1.48	0.219
Residual	36	15179354.	421649.		

So basically the linear slope for N is the same for the two planting methods ($P=0.214$), however there is a difference in the non-linear trend (this is marked as Deviations in the Nitrogen.Planting_Method interaction. Deviations just represents all other terms that have not been modelled. With 3 nitrogen rates there can only be a linear and a quadratic component, since a quadratic equation will pass through three points perfectly. Here we fitted just the linear term – via $POL(Nitrogen;1)$.) The difference in the quadratic component is visible in the previous means plot.

4. There is a strongly significant difference in the means for broadcasting versus transplanting across the varieties ($P < 0.001$), as seen in the means plot:



5. There is also a significant difference in response of the varieties to nitrogen ($P = 0.034$). It is clear that not all slopes are the same ($P = 0.005$), while there is no evidence of any particular quadratic component overall ($P = 0.974$), or across varieties ($P = 0.865$):



6. One of the drawbacks of these stratified designs is that simple comparisons of means are, in many cases, not exactly t tests. This can be picked up from the s.e.d. and l.s.d. values. Any comparison that shows in GenStat's output with degrees of freedom that is not an integer results from a formula in a Satterthwaite-like approximate t test. In fact, for this design, few comparisons between two-way means are exact t tests, so care needs to be used when selecting the correct s.e.d. value to use for a particular comparison. For example, to test the mean difference between broadcast and transplant methods across the varieties, you would use 414.6 on 13.31 df for each of these differences (because the same varieties is used for each difference):

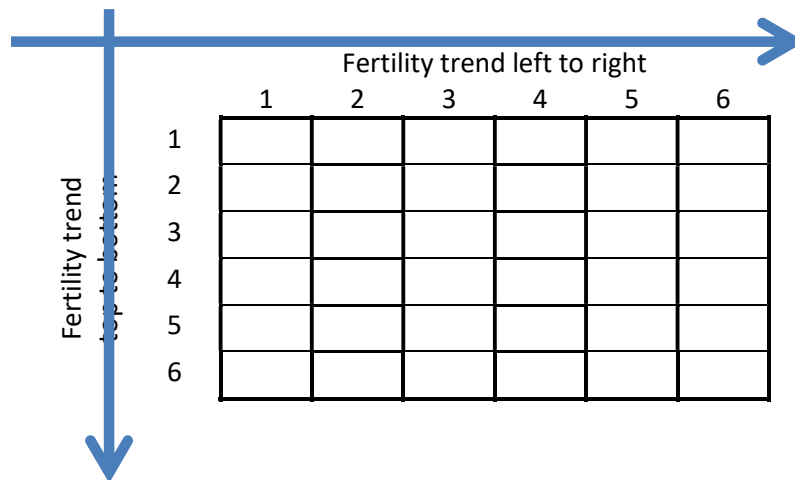
Variety	Planting Method		difference
	Broadcast	Transplanted	
IR8	5417	4898	519
IR127-8-1-10	6286	5540	746
IR305-4-12-1-13	6080	6097	-17
IR400-2-5-3-3-2	5569	6198	-629
IR665-58	5249	4839	410
Peta	3138	5150	-2012

However, to compare two varietal means that were both broadcast (or transplanted) the t test is approximate, the s.e.d. value to use is 414.6 and there are only 13.31 degrees of freedom available for the test

Latin Square design

The next design is useful when there are *two fertility gradients* in the field. In a way, the LS design consists of a single replicate of the strip-plot design just considered.

Consider the following layout:



Suppose the trends are such that plots become drier from left to right and top to bottom, so that the top left hand area is a higher yielding area, declining towards the bottom right hand corner. Now suppose we had randomized 6 replicates of each of 6 treatments (T11 to T6) into the field and obtained something like:

		Fertility trend left to right					
		1	2	3	4	5	6
Fertility trend top to bottom	1	T1	T6	T3	T6	T5	T5
	2	T5	T1	T1	T4	T3	T4
	3	T6	T1	T4	T6	T2	T2
	4	T4	T6	T6	T1	T4	T2
	5	T1	T5	T3	T3	T2	T3
	6	T3	T5	T4	T2	T5	T2

A comparison between the mean for T1 and for T6 would be “fair”, because all 6 replicates of each treatment were allocated roughly equally to the higher yielding area of the field. Not so a comparison between the mean for T1 or for T6 compared to T2 which was (randomly) allocated mainly in the lower yielding plots. So, a difference between the mean of T1 and of T2 would not reflect the true treatment difference only, but would contain a component of the difference in yield due to the higher versus lower yielding areas.

The only way to be fair for every treatment comparison is to balance the layout so that *every treatment occurs once in each row and once in each column*, so something like:

		Fertility trend left to right					
		1	2	3	4	5	6
Fertility trend top to bottom	1	T4	T5	T1	T2	T3	T6
	2	T3	T1	T6	T4	T5	T2
	3	T5	T6	T2	T3	T1	T4
	4	T1	T2	T4	T5	T6	T3
	5	T6	T4	T3	T1	T2	T5
	6	T2	T3	T5	T6	T4	T1

Now when you compare the means of T1 and T2, each mean involves the average row fertility effect, as well as the average column fertility effect. Consequently, when you calculate the *mean difference*, the average row fertility effect and the average column fertility effect both disappear from the difference, leaving behind an estimate of the real difference in means.

A design balanced in this way is known mathematically as a **Latin Square design**. Note that to achieve such a balance, there must be the same number of replicates as there are treatments.

So before looking at experimental data, let's examine what we do in the field for t treatments.

 We construct t rows in the field, blocking (allowing) for the trend in that direction. This gives rise to a *row block stratum*, and, as for an RCB ANOVA, will have the **Row** block component alone.

 We next construct t columns in the field, across the rows, blocking for the trend in that direction. This gives rise to a *column block stratum*, and, as for an RCB ANOVA, will have the **Column** block component alone.

 The intersection of the row and column blocks produces t^2 plots in the field. Mathematicians have shown us how to achieve the appropriate balance of treatments so that each treatment occurs once in each row and once in each column. But these individual plots are third stratum and allow the treatments to be tested.

The Block Structure in GenStat comes from what was done in the field:

Construct row blocks $\Rightarrow$ **Row** stratum

Construct column blocks $\Rightarrow$ **Column** stratum

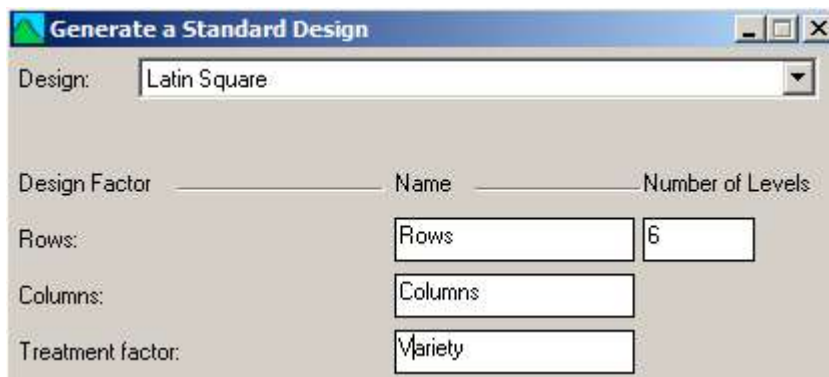
Plots at the intersection of row blocks and column blocks $\Rightarrow$ **Row.Column** stratum

Combining the 3 strata: **Row+Column+Row.Column** $\Leftarrow$ **Row*Column** using GenStat's rules.

The Latin Square ANOVA would therefore appear like this:

Component	df
Row stratum	
Row	$(t-1)$
Column stratum	
Column	$(t-1)$
Row.Column stratum	
Treatment	$(t-1)$
Residual	$(t-1)(t-2)$
Total	t^2-1

GenStat has a menu to produce randomized designs, and even show what the analysis will look like. For this use Stats > Design > Generate a Standard Design. Select Latin Square from the drop down list, name treatments and indicate their number. The design is often used for animal trials, where a number of animals are treated but the set of treatments is re-applied to the animals at different periods of time. Then you would name the rows and columns Animal and Period say.



Generate a Standard Design

Design: Latin Square

Design Factor	Name	Number of Levels
Rows:	Rows	6
Columns:	Columns	
Treatment factor:	Variety	

Grain yield of three promising maize hybrids (A, B, D) and a check hybrid (C) set out in a 4×4 Latin Square design, from Page 33 of G&G

Design Column					Yield (t/ha) Column				Row means
	1	2	3	4	1	2	3	4	
1	B	D	C	A	1.640	1.210	1.425	1.345	1.405
2	C	A	D	B	1.475	1.185	1.400	1.290	1.338
3	A	C	B	D	1.670	0.710	1.665	1.180	1.306
4	D	B	A	C	1.565	1.290	1.655	0.660	1.293
Column means					1.588	1.099	1.536	1.119	

In fact, there doesn't appear to be much difference in the row means, the variance for which is 0.003. Each is a mean of 4 plot yields, so the **Row m.s.** in the ANOVA would be $4 \times 0.003 = 0.010$. (Note we are rounding off to 3 dp but precision occurs in the unseen decimals.)

Column means are somewhat more varied, their variance is 0.069 and hence the **Column m.s.** in the ANOVA would be $4 \times 0.069 = 0.276$.

To obtain **Hybrid m.s.** we need their means, which entails averaging the plot yields that received Hybrid A, etc:

A	B	C	D
1.464	1.471	1.068	1.339

The variance of these 4 means is 0.036, and so the **Hybrid m.s.** in the ANOVA would be $4 \times 0.036 = 0.142$.

The Residual m.s. in the ANOVA would be the variance of the $t^2 = 16$ residuals, using $(t-1)(t-2) = 6$ as the divisor.

Analysis of variance

Source of variation	d.f.	s.s.	m.s.	v.r.	F pr.
Row stratum	3	0.03015	0.01005	0.47	
Column stratum	3	0.82734	0.27578	12.77	
Row.Column stratum					
Variety	3	0.42684	0.14228	6.59	0.025
Residual	6	0.12958	0.02160		
Total	15	1.41392			

Tables of means

Variate: Yield

Grand mean 1.335

Variety	A	B	C	D
	1.464	1.471	1.068	1.339

Standard errors of differences of means

Table	Variety
rep.	4
d.f.	6
s.e.d.	0.1039

Least significant differences of means (5% level)

Table	Variety
rep.	4
d.f.	6
l.s.d.	0.2543

None of the three promising maize hybrids A, B or D differs from another (the 5% l.s.d. value is 0.254), however all three are significantly different (at 5% at least) from the check maize variety

D. Had you set up 3 simple comparisons as part of the ANOVA you would have the stronger statistical evidence. The P values for comparisons of each of the three hybrids with C are:

A vs C (P=0.009), B vs C (P=0.008) and D vs C (P=0.040).

FURTHER READING

For an introductory, practical and illustrative guide to design of experiments and data analysis we refer you to following book:

Welham, S.J., S.A Gezan, S.J. Clark and A. Mead. 2015. Statistical Methods in Biology: Design and Analysis of Experiments and Regression. CRC Press, Boca Rotan, FL.

For introductory information on ANOVA & REML methods we refer you to the following manual:

O'Neill, M. 2010. ANOVA & REML: A Guide to Linear Mixed Models in an Experimental Design Context. Statistical Advisory & Training Services Pty Ltd. NSW, Australia. www.stats.net.au

OTHER REFERENCES:

Bowley, S.R. 2015. A Hitchhiker's Guide to Statistics in Plant Biology, GLMM Edition. Plants et al. Inc., Guelph Canada.

Gbur, E.E., W.W. Stroup, K.S. McCarter, S. Durham, L.J. Young, M. Christman, M. West and M. Kramer. 2012. Generalized Linear Mixed Models in the Agricultural and Natural Resources Sciences. Madison, WI: Agronomy Society of America.

Gomez, K.A, and A.A. Gomez. 1984. Statistical Procedures for Agriculture Research. 2nd Ed. John Wiley and Sons. New York.

Lee, C.J., M. O'Donnell and M. O'Neill. 2008. Statistical Analysis of Field Trials with Changing Treatment Variance. Agron. J. 100:484–489.

Mead, R., S.G. Gilmour and A. Mead. 2012. Statistical Principles for the Design of Experiments. Cambridge University Press. Cambridge, UK.

Payne, R.W. 2013. The Guide to the GenStat® Command Language (Release 16) Part 2: Statistics. VSN International, 5 The Waterhouse, Waterhouse Street, Hemel Hempstead, Hertfordshire HP1 1ES, UK

Payne, R.W., D. Murray, S. Harding, D. Baird and D. Souter. 2013. Introduction to GenStat for Windows. VSN International, 5 The Waterhouse, Waterhouse Street, Hemel Hempstead, Hertfordshire HP1 1ES, UK

Peterson, R.G. 1994. Agricultural Field Experiments: Design and Analysis. Marcel Dekker, Inc., New York, NY.

Pierce, S.C., G.M. Clark, G.V. Dyke, R.E. Kempson. 1988. Manual of Crop Experimentation. Oxford University Press, New York.

Senn, S. 2003. A Conversation with John Nelder. Statistical Science. Vol. 18, No. 1. 118-131. Institute of Mathematical Science.

Stroup, W.W. 2012. Generalized Linear Mixed Models: Modern Concepts, Methods and Applications. CRC Press. Boca Raton, FL.

Stroup, W.W. 2013. Rethinking the Analysis of Non-Normal Data in Plant and Soil Science Agron J., Madison, WI: Agronomy Society of America.

VSN International Ltd, 5 The Waterhouse, Waterhouse Street Hemel Hempstead, *HP1 1ES UK*.
www.vsnl.co.uk

Williams, E.R., A.C Matheson and C.E. Harwood. 2002. Experimental Design and Analysis for Tree Improvement. 2nd Edition. CSIRCO Publishing, Australia

Appendix A. Data sets from Statistical Procedures for Agricultural Research, 2nd Edition, Kwanchai A. Gomez and Arturo A. Gomez.

Copyright © 1984 by John Wiley & Sons, Inc.

CRD - page 14

Treatment	Rep	Yield
Dol-Mix (1 kg)	1	2537
Dol-Mix (1 kg)	2	2069
Dol-Mix (1 kg)	3	2104
Dol-Mix (1 kg)	4	1797
Dol-Mix (2 kg)	1	3366
Dol-Mix (2 kg)	2	2591
Dol-Mix (2 kg)	3	2211
Dol-Mix (2 kg)	4	2544
DDT + γ -BHC	1	2536
DDT + γ -BHC	2	2459
DDT + γ -BHC	3	2827
DDT + γ -BHC	4	2385
Azodrin	1	2387
Azodrin	2	2453
Azodrin	3	1556
Azodrin	4	2116
Dicecron-Boom	1	1997
Dicecron-Boom	2	1679
Dicecron-Boom	3	1649
Dicecron-Boom	4	1859
Dicecron-Knap	1	1796
Dicecron-Knap	2	1704
Dicecron-Knap	3	1904
Dicecron-Knap	4	1320
Control	1	1401
Control	2	1516
Control	3	1270
Control	4	1077

Treatment	Rep	Yield
25	1	5113
50	1	5346
75	1	5272
100	1	5164
125	1	4804
150	1	5254
25	2	5398
50	2	5952
75	2	5713
100	2	4831
125	2	4848
150	2	4542
25	3	5307
50	3	4719
75	3	5483
100	3	4986
125	3	4432
150	3	4919
25	4	4678
50	4	4264
75	4	4749
100	4	4410
125	4	4748
150	4	4098

Latin square - page 33

Row	Column	Variety	Yield
1	1	B	1.64
1	2	D	1.21
1	3	C	1.425
1	4	A	1.345
2	1	C	1.475
2	2	A	1.185
2	3	D	1.4
2	4	B	1.29
3	1	A	1.67
3	2	C	0.71
3	3	B	1.665
3	4	D	1.18
4	1	D	1.565
4	2	B	1.29
4	3	A	1.655
4	4	C	0.66

Split-plot - page 102

Main plot = Nitrogen, Subplot = Variety, Replicates = 3

Variety	Nitrogen	Rep	Yield
IR8	0	1	4430
IR5	0	1	3944
C4-63	0	1	3464
Peta	0	1	4126
IR8	0	2	4478
IR5	0	2	5314
C4-63	0	2	2944
Peta	0	2	4482
IR8	0	3	3850
IR5	0	3	3660
C4-63	0	3	3142
Peta	0	3	4836
IR8	60	1	5418
IR5	60	1	6502
C4-63	60	1	4768
Peta	60	1	5192
IR8	60	2	5166
IR5	60	2	5858
C4-63	60	2	6004
Peta	60	2	4604
IR8	60	3	6432
IR5	60	3	5586
C4-63	60	3	5556
Peta	60	3	4652
IR8	90	1	6076
IR5	90	1	6008
C4-63	90	1	6244
Peta	90	1	4546
IR8	90	2	6420
IR5	90	2	6127
C4-63	90	2	5724
Peta	90	2	5744
IR8	90	3	6704
IR5	90	3	6642
C4-63	90	3	6014
Peta	90	3	4146
IR8	120	1	6462
IR5	120	1	7139
C4-63	120	1	5792
Peta	120	1	2774

IR8	120	2	7056
IR5	120	2	6982
C4-63	120	2	5880
Peta	120	2	5036
IR8	120	3	6680
IR5	120	3	6564
C4-63	120	3	6370
Peta	120	3	3638
IR8	150	1	7290
IR5	150	1	7682
C4-63	150	1	7080
Peta	150	1	1414
IR8	150	2	7848
IR5	150	2	6594
C4-63	150	2	6662
Peta	150	2	1960
IR8	150	3	7552
IR5	150	3	6576
C4-63	150	3	6320
Peta	150	3	2766
IR8	180	1	8452
IR5	180	1	6228
C4-63	180	1	5594
Peta	180	1	2248
IR8	180	2	8832
IR5	180	2	7387
C4-63	180	2	7122
Peta	180	2	1380
IR8	180	3	8818
IR5	180	3	6006
C4-63	180	3	5480
Peta	180	3	2014

Strip-plot - page 110

Horizontal factor = Nitrogen, Vertical factor =
Variety

Nitrogen	Variety	Rep	Yield
0	IR8	1	2373
60	IR8	1	4076
120	IR8	1	7254
0	IR8	2	3958
60	IR8	2	6431
120	IR8	2	6808
0	IR8	3	4384
60	IR8	3	4889
120	IR8	3	8582
0	IR127-80	1	4007
60	IR127-80	1	5630
120	IR127-80	1	7053
0	IR127-80	2	5795
60	IR127-80	2	7334
120	IR127-80	2	8284
0	IR127-80	3	5001
60	IR127-80	3	7177
120	IR127-80	3	6297
0	IR305-4-12	1	2620
60	IR305-4-12	1	4676
120	IR305-4-12	1	7666
0	IR305-4-12	2	4508
60	IR305-4-12	2	6672
120	IR305-4-12	2	7328
0	IR305-4-12	3	5621
60	IR305-4-12	3	7019
120	IR305-4-12	3	8611
0	IR400-2-5	1	2726
60	IR400-2-5	1	4838
120	IR400-2-5	1	6881
0	IR400-2-5	2	5630
60	IR400-2-5	2	7007
120	IR400-2-5	2	7735
0	IR400-2-5	3	3821
60	IR400-2-5	3	4816
120	IR400-2-5	3	6667
0	IR665-58	1	4447
60	IR665-58	1	5549
120	IR665-58	1	6880

0	IR665-58	2	3276
60	IR665-58	2	5340
120	IR665-58	2	5080
0	IR665-58	3	4582
60	IR665-58	3	6011
120	IR665-58	3	6076
0	Peta	1	2572
60	Peta	1	3896
120	Peta	1	1556
0	Peta	2	3724
60	Peta	2	2822
120	Peta	2	2706
0	Peta	3	3326
60	Peta	3	4425
120	Peta	3	3214

Split-split-plot - page 143

Main plot = Nitrogen, Subplot = Management, Sub-subplot = Variety

Management	Variety	Nitrogen	Rep	Yield
Minimum	V1	0	1	3.32
Minimum	V1	0	2	3.864
Minimum	V1	0	3	4.507
Minimum	V2	0	1	6.101
Minimum	V2	0	2	5.122
Minimum	V2	0	3	4.815
Minimum	V3	0	1	5.355
Minimum	V3	0	2	5.536
Minimum	V3	0	3	5.244
Optimum	V1	0	1	3.766
Optimum	V1	0	2	4.311
Optimum	V1	0	3	4.875
Optimum	V2	0	1	5.096
Optimum	V2	0	2	4.873
Optimum	V2	0	3	4.166
Optimum	V3	0	1	7.442
Optimum	V3	0	2	6.462
Optimum	V3	0	3	5.584
Intensive	V1	0	1	4.66
Intensive	V1	0	2	5.915
Intensive	V1	0	3	5.4
Intensive	V2	0	1	6.573
Intensive	V2	0	2	5.495
Intensive	V2	0	3	4.225
Intensive	V3	0	1	7.018
Intensive	V3	0	2	8.02
Intensive	V3	0	3	7.642
Minimum	V1	50	1	3.188
Minimum	V1	50	2	4.752
Minimum	V1	50	3	4.756
Minimum	V2	50	1	5.595
Minimum	V2	50	2	6.78
Minimum	V2	50	3	5.39
Minimum	V3	50	1	6.706
Minimum	V3	50	2	6.546
Minimum	V3	50	3	7.092
Optimum	V1	50	1	3.625
Optimum	V1	50	2	4.809
Optimum	V1	50	3	5.295
Optimum	V2	50	1	6.357

Optimum	V2	50	2	5.925
Optimum	V2	50	3	5.163
Optimum	V3	50	1	8.592
Optimum	V3	50	2	7.646
Optimum	V3	50	3	7.212
Intensive	V1	50	1	5.232
Intensive	V1	50	2	5.17
Intensive	V1	50	3	6.046
Intensive	V2	50	1	7.016
Intensive	V2	50	2	7.442
Intensive	V2	50	3	4.478
Intensive	V3	50	1	8.48
Intensive	V3	50	2	9.942
Intensive	V3	50	3	8.714
Minimum	V1	80	1	5.468
Minimum	V1	80	2	5.788
Minimum	V1	80	3	4.422
Minimum	V2	80	1	5.442
Minimum	V2	80	2	5.988
Minimum	V2	80	3	6.509
Minimum	V3	80	1	8.452
Minimum	V3	80	2	6.698
Minimum	V3	80	3	8.65
Optimum	V1	80	1	5.759
Optimum	V1	80	2	6.13
Optimum	V1	80	3	5.308
Optimum	V2	80	1	6.398
Optimum	V2	80	2	6.533
Optimum	V2	80	3	6.569
Optimum	V3	80	1	8.662
Optimum	V3	80	2	8.526
Optimum	V3	80	3	8.514
Intensive	V1	80	1	6.215
Intensive	V1	80	2	7.106
Intensive	V1	80	3	6.318
Intensive	V2	80	1	6.953
Intensive	V2	80	2	6.914
Intensive	V2	80	3	7.991
Intensive	V3	80	1	9.112
Intensive	V3	80	2	9.14
Intensive	V3	80	3	9.32
Minimum	V1	110	1	4.246
Minimum	V1	110	2	4.842

Minimum	V1	110	3	4.863
Minimum	V2	110	1	6.209
Minimum	V2	110	2	6.768
Minimum	V2	110	3	5.779
Minimum	V3	110	1	8.042
Minimum	V3	110	2	7.414
Minimum	V3	110	3	6.902
Optimum	V1	110	1	5.255
Optimum	V1	110	2	5.742
Optimum	V1	110	3	5.345
Optimum	V2	110	1	6.992
Optimum	V2	110	2	7.856
Optimum	V2	110	3	6.164
Optimum	V3	110	1	9.08
Optimum	V3	110	2	9.016
Optimum	V3	110	3	7.778
Intensive	V1	110	1	6.829
Intensive	V1	110	2	5.869
Intensive	V1	110	3	6.011
Intensive	V2	110	1	7.565
Intensive	V2	110	2	7.626
Intensive	V2	110	3	7.362
Intensive	V3	110	1	9.66
Intensive	V3	110	2	8.966
Intensive	V3	110	3	9.128
Minimum	V1	140	1	3.132
Minimum	V1	140	2	4.375
Minimum	V1	140	3	4.678
Minimum	V2	140	1	6.86
Minimum	V2	140	2	6.894
Minimum	V2	140	3	6.573
Minimum	V3	140	1	9.314
Minimum	V3	140	2	8.508
Minimum	V3	140	3	8.032
Optimum	V1	140	1	5.389
Optimum	V1	140	2	4.315
Optimum	V1	140	3	5.896
Optimum	V2	140	1	6.857
Optimum	V2	140	2	6.974
Optimum	V2	140	3	7.422
Optimum	V3	140	1	9.224
Optimum	V3	140	2	9.68
Optimum	V3	140	3	9.294

Intensive	V1	140	1	5.217
Intensive	V1	140	2	5.389
Intensive	V1	140	3	7.309
Intensive	V2	140	1	7.254
Intensive	V2	140	2	7.812
Intensive	V2	140	3	8.95
Intensive	V3	140	1	10.36
Intensive	V3	140	2	9.896
Intensive	V3	140	3	9.712

Strip-split-plot - page 155

Horizontal strip = Variety, Vertical strip = Nitrogen, Subplot = Planting method

Variety	Nitrogen	Planting Method	Rep	Yield
IR8	0	Broadcast	1	2373
IR8	0	Broadcast	2	3958
IR8	0	Broadcast	3	4384
IR8	0	Transplanted	1	2293
IR8	0	Transplanted	2	3528
IR8	0	Transplanted	3	2538
IR127-8-1-10	0	Broadcast	1	4007
IR127-8-1-10	0	Broadcast	2	5795
IR127-8-1-10	0	Broadcast	3	5001
IR127-8-1-10	0	Transplanted	1	4035
IR127-8-1-10	0	Transplanted	2	4885
IR127-8-1-10	0	Transplanted	3	4583
IR305-4-12-1-3	0	Broadcast	1	2620
IR305-4-12-1-3	0	Broadcast	2	4508
IR305-4-12-1-3	0	Broadcast	3	5621
IR305-4-12-1-3	0	Transplanted	1	4527
IR305-4-12-1-3	0	Transplanted	2	4866
IR305-4-12-1-3	0	Transplanted	3	3628
IR400-2-5-3-3-2	0	Broadcast	1	2726
IR400-2-5-3-3-2	0	Broadcast	2	5630
IR400-2-5-3-3-2	0	Broadcast	3	3821
IR400-2-5-3-3-2	0	Transplanted	1	5274
IR400-2-5-3-3-2	0	Transplanted	2	6200
IR400-2-5-3-3-2	0	Transplanted	3	4038
IR665-58	0	Broadcast	1	4447
IR665-58	0	Broadcast	2	3276
IR665-58	0	Broadcast	3	4582
IR665-58	0	Transplanted	1	4655
IR665-58	0	Transplanted	2	2796
IR665-58	0	Transplanted	3	3739
Peta	0	Broadcast	1	2572
Peta	0	Broadcast	2	3724
Peta	0	Broadcast	3	3326
Peta	0	Transplanted	1	4535
Peta	0	Transplanted	2	5457
Peta	0	Transplanted	3	3537
IR8	60	Broadcast	1	4076
IR8	60	Broadcast	2	6431

IR8	60	Broadcast	3	4889
IR8	60	Transplanted	1	3085
IR8	60	Transplanted	2	7502
IR8	60	Transplanted	3	4362
IR127-8-1-10	60	Broadcast	1	5630
IR127-8-1-10	60	Broadcast	2	7334
IR127-8-1-10	60	Broadcast	3	7177
IR127-8-1-10	60	Transplanted	1	3728
IR127-8-1-10	60	Transplanted	2	7424
IR127-8-1-10	60	Transplanted	3	5377
IR305-4-12-1-3	60	Broadcast	1	4676
IR305-4-12-1-3	60	Broadcast	2	6672
IR305-4-12-1-3	60	Broadcast	3	7019
IR305-4-12-1-3	60	Transplanted	1	4946
IR305-4-12-1-3	60	Transplanted	2	7611
IR305-4-12-1-3	60	Transplanted	3	6142
IR400-2-5-3-3-2	60	Broadcast	1	4838
IR400-2-5-3-3-2	60	Broadcast	2	7007
IR400-2-5-3-3-2	60	Broadcast	3	4816
IR400-2-5-3-3-2	60	Transplanted	1	4878
IR400-2-5-3-3-2	60	Transplanted	2	6928
IR400-2-5-3-3-2	60	Transplanted	3	4829
IR665-58	60	Broadcast	1	5549
IR665-58	60	Broadcast	2	5340
IR665-58	60	Broadcast	3	6011
IR665-58	60	Transplanted	1	4646
IR665-58	60	Transplanted	2	5006
IR665-58	60	Transplanted	3	4666
Peta	60	Broadcast	1	3896
Peta	60	Broadcast	2	2822
Peta	60	Broadcast	3	4425
Peta	60	Transplanted	1	4627
Peta	60	Transplanted	2	4461
Peta	60	Transplanted	3	4774
IR8	120	Broadcast	1	7254
IR8	120	Broadcast	2	6808
IR8	120	Broadcast	3	8582
IR8	120	Transplanted	1	6661
IR8	120	Transplanted	2	6353
IR8	120	Transplanted	3	7759
IR127-8-1-10	120	Broadcast	1	7053
IR127-8-1-10	120	Broadcast	2	8284
IR127-8-1-10	120	Broadcast	3	6297

IR127-8-1-10	120	Transplanted	1	6440
IR127-8-1-10	120	Transplanted	2	7648
IR127-8-1-10	120	Transplanted	3	5736
IR305-4-12-1-3	120	Broadcast	1	7666
IR305-4-12-1-3	120	Broadcast	2	7328
IR305-4-12-1-3	120	Broadcast	3	8611
IR305-4-12-1-3	120	Transplanted	1	8632
IR305-4-12-1-3	120	Transplanted	2	7101
IR305-4-12-1-3	120	Transplanted	3	7416
IR400-2-5-3-3-2	120	Broadcast	1	6881
IR400-2-5-3-3-2	120	Broadcast	2	7735
IR400-2-5-3-3-2	120	Broadcast	3	6667
IR400-2-5-3-3-2	120	Transplanted	1	6545
IR400-2-5-3-3-2	120	Transplanted	2	9838
IR400-2-5-3-3-2	120	Transplanted	3	7253
IR665-58	120	Broadcast	1	6880
IR665-58	120	Broadcast	2	5080
IR665-58	120	Broadcast	3	6076
IR665-58	120	Transplanted	1	6995
IR665-58	120	Transplanted	2	4486
IR665-58	120	Transplanted	3	6564
Peta	120	Broadcast	1	1556
Peta	120	Broadcast	2	2706
Peta	120	Broadcast	3	3214
Peta	120	Transplanted	1	5374
Peta	120	Transplanted	2	7218
Peta	120	Transplanted	3	6369

Mick O'Neill
Statistical Advisory & Training Service Pty Ltd.
21 Days Crescent, Blackheath NWS Australia
www.stats.net.au



Curt Lee
Agro-Tech, Inc.
4489 Highway 41 North, Velva North Dakota USA
www.agrotechresearch.com

